

Lectures for the 27th IAU ISYA
Ifrane, 2nd - 23rd July 2004



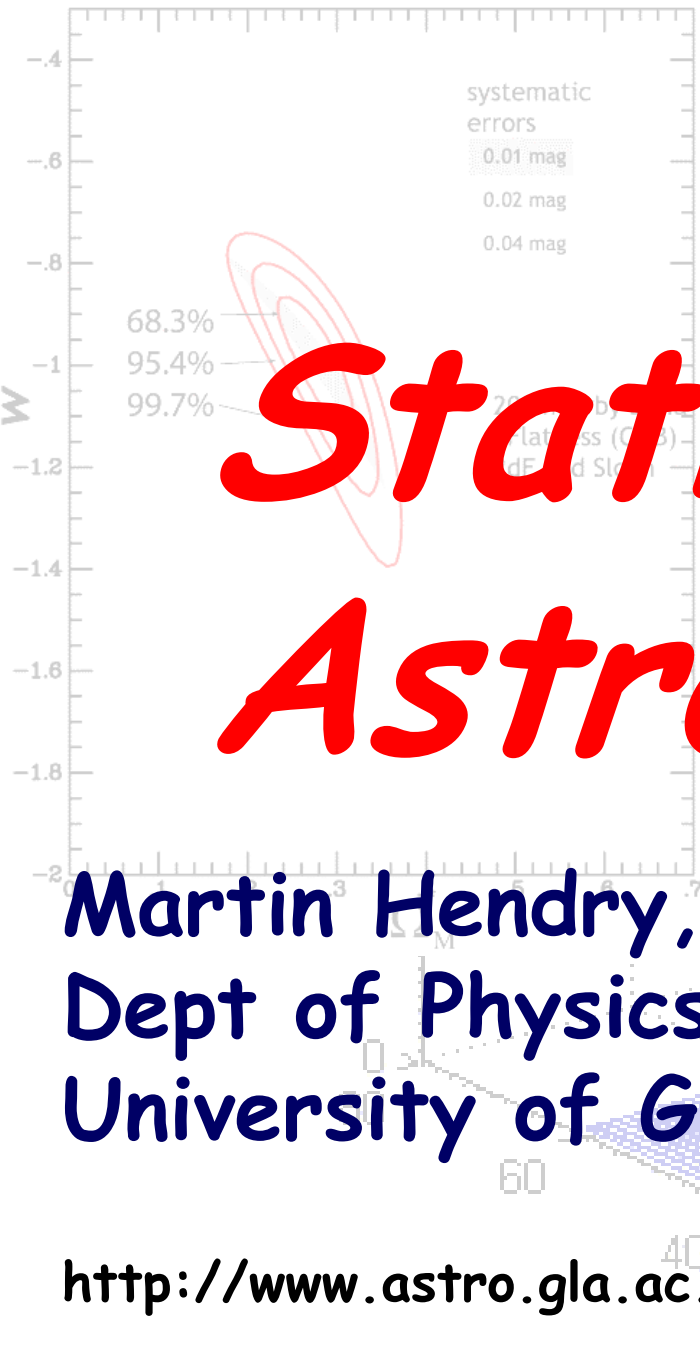
UNIVERSITY
of
GLASGOW



Statistical Astronomy

Martin Hendry,
Dept of Physics and Astronomy
University of Glasgow, UK

<http://www.astro.gla.ac.uk/users/martin/isya/>



$$p(x | y, I) = \frac{p(y | x, I) p(x, I)}{p(y, I)}$$

Parameter estimation: the Gaussian approximation

Posterior

Likelihood

Prior

$$p(\text{model} \mid \text{data}, I) \propto p(\text{data} \mid \text{model}, I) \times p(\text{model} \mid I)$$



Parameter estimation: the Gaussian approximation

$$p(\theta \mid \text{data}, I) \propto p(\text{data} \mid \theta, I) \times p(\theta \mid I)$$

'Best' estimator: $\left. \frac{\partial p(\theta \mid \text{data}, I)}{\partial \theta} \right|_{\theta=\theta_0} = 0$ ← Maximise posterior likelihood

Equivalently, we can define $\ell = \log p(\theta \mid \text{data}, I)$ and compute $\left. \frac{\partial \ell}{\partial \theta} \right|_{\theta=\theta_0} = 0$

Taylor expand $\ell(\theta)$ around $\theta=\theta_0$:

$$\ell(\theta) = \ell(\theta_0) + \cancel{\left. \frac{\partial \ell}{\partial \theta} \right|_{\theta=\theta_0} (\theta - \theta_0)} + \frac{1}{2} \left. \frac{\partial^2 \ell}{\partial \theta^2} \right|_{\theta=\theta_0} (\theta - \theta_0)^2 + \dots$$



Parameter estimation: the Gaussian approximation

$$p(\theta \mid \text{data}, I) = \exp [\ell(\theta)]$$

Neglecting higher order terms in $\ell(\theta)$ ← Gaussian approximation

$$p(\theta \mid \text{data}, I) \propto \exp \left(-\frac{A}{2} (\theta - \theta_0)^2 \right)$$

where $A = -\frac{\partial^2 \ell}{\partial \theta^2} \bigg|_{\theta=\theta_0}$

This is equivalent to a normal distribution, with $\sigma^{-2} = A = -\frac{\partial^2 \ell}{\partial \theta^2} \bigg|_{\theta=\theta_0}$

Can summarise inference from posterior by

$$\theta = \theta_0 \pm \sigma$$



Parameter estimation: 2-D case

Recall our definition of *variance*

$$\text{var}[x] = \int_{-\infty}^{\infty} (x - \langle x \rangle)^2 p(x | I) dx$$

Extend to 2 variables - *covariance*

$$\text{cov}[x, y] = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (x - \langle x \rangle)(y - \langle y \rangle) p(x, y | I) dx dy$$

If x and y are *independent*, $\text{cov}[x, y] = 0$

This is because $p(x, y | I) = p(x | I)p(y | I)$

Parameter estimation: 2-D case

$$p(\theta_1, \theta_2 \mid \text{data}, I) \propto p(\text{data} \mid \theta_1, \theta_2, I) \times p(\theta_1, \theta_2 \mid I)$$

'Best' estimator: $\left. \frac{\partial p(\theta_1, \theta_2 \mid \text{data}, I)}{\partial \theta_j} \right|_{\theta_j = \theta_{0j}} = 0$

Compute $\left. \frac{\partial \ell}{\partial \theta_j} \right|_{\theta_j = \theta_{0j}} = 0$ where $\ell = \log p(\theta_1, \theta_2 \mid \text{data}, I)$



Parameter estimation: 2-D case

Taylor expand $\ell(\theta_1, \theta_2)$ around θ_{0j} :

$$\begin{aligned} \ell(\theta_1, \theta_2) = & \ell(\theta_{01}, \theta_{02}) + \frac{\partial \ell}{\partial \theta_1} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01}) + \frac{\partial \ell}{\partial \theta_2} \bigg|_{\theta_j = \theta_{0j}} (\theta_2 - \theta_{02}) + \\ & \frac{1}{2} \left[\frac{\partial^2 \ell}{\partial \theta_1^2} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01})^2 + \frac{\partial^2 \ell}{\partial \theta_2^2} \bigg|_{\theta_j = \theta_{0j}} (\theta_2 - \theta_{02})^2 + 2 \frac{\partial^2 \ell}{\partial \theta_1 \partial \theta_2} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01})(\theta_2 - \theta_{02}) \right] + \dots \end{aligned}$$

$$p(\theta_1, \theta_2 \mid \text{data}, I) \propto \exp [\ell(\theta_1, \theta_2)]$$

$$\propto \exp \left[-\frac{1}{2} Q \right] \quad \leftarrow \text{Gaussian approximation}$$

$$\chi^2 \nearrow$$

Maximising likelihood
 $\equiv$ Minimising χ^2



Parameter estimation: 2-D case

Taylor expand $\ell(\theta_1, \theta_2)$ around θ_{0j} :

$$Q = (\theta_1 - \theta_{10} \quad \theta_2 - \theta_{20}) \begin{bmatrix} A & C \\ C & B \end{bmatrix} \begin{pmatrix} \theta_1 - \theta_{10} \\ \theta_2 - \theta_{20} \end{pmatrix}$$

Quadratic
form

where

$$A = \frac{\partial^2 \ell}{\partial \theta_1^2} \bigg|_{\theta_j = \theta_{0j}} \quad B = \frac{\partial^2 \ell}{\partial \theta_2^2} \bigg|_{\theta_j = \theta_{0j}} \quad C = \frac{\partial^2 \ell}{\partial \theta_1 \partial \theta_2} \bigg|_{\theta_j = \theta_{0j}}$$

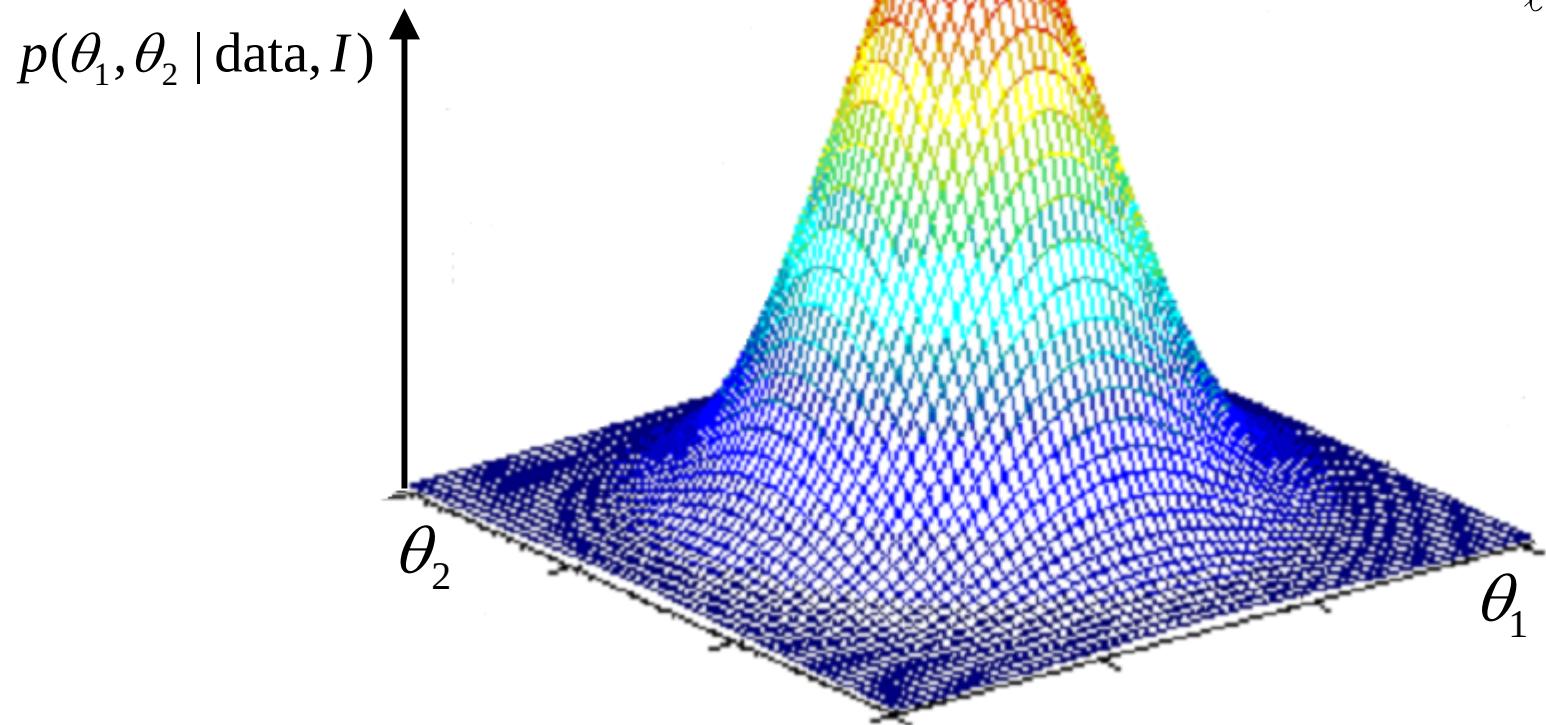
This is a bivariate normal distribution with covariance matrix

$$\sigma_{ij}^2 = \text{cov}_{ij} = \langle (\theta_i - \theta_{i0})(\theta_j - \theta_{j0}) \rangle = \left[-\frac{\partial^2 \ell}{\partial \theta_i \partial \theta_j} \right]^{-1}$$

Fisher information matrix



Parameter estimation: 2-D case



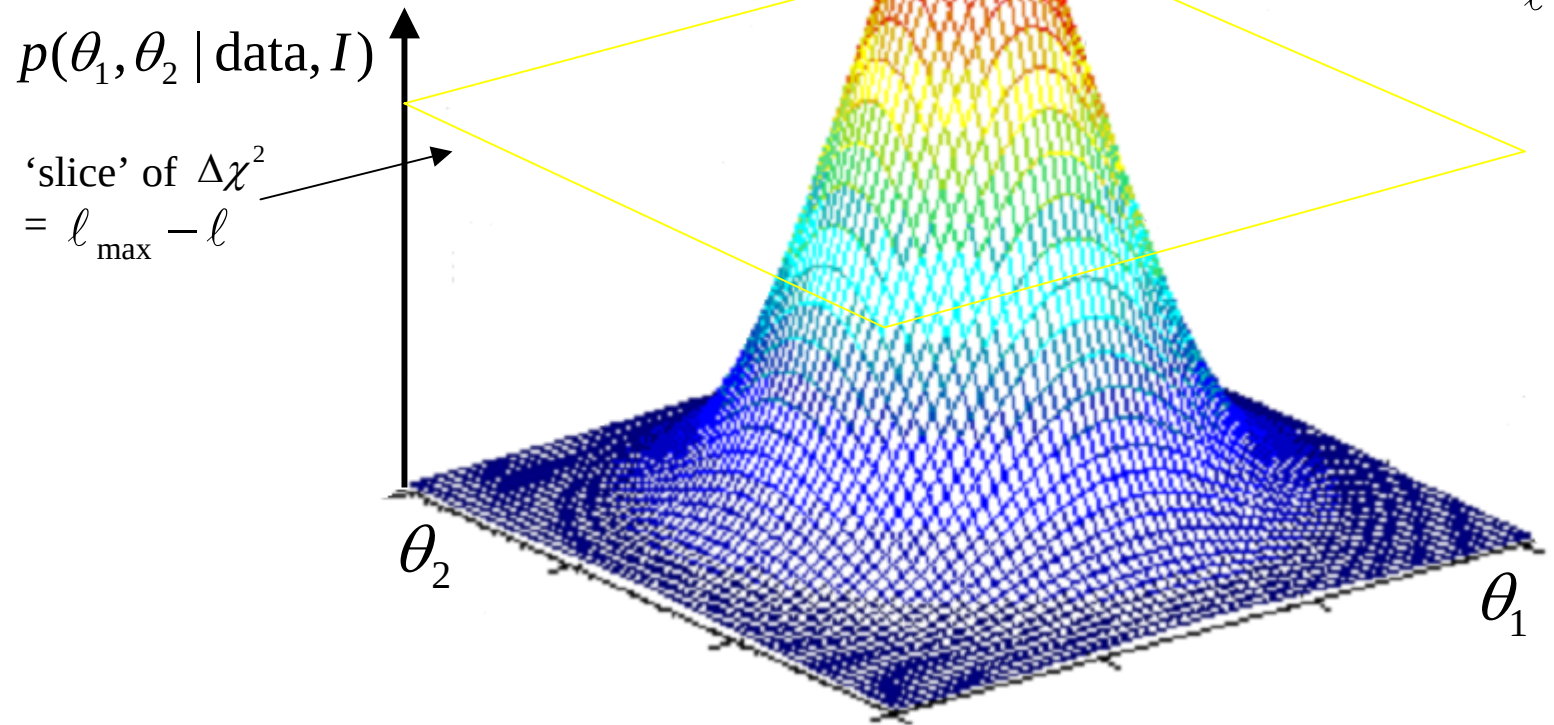
This is a bivariate normal distribution with covariance matrix

$$\sigma_{ij}^2 = \text{cov}_{ij} = \langle (\theta_i - \theta_{i0})(\theta_j - \theta_{j0}) \rangle = \left[-\frac{\partial^2 \ell}{\partial \theta_i \partial \theta_j} \right]^{-1}$$

Fisher information matrix



Parameter estimation: 2-D case



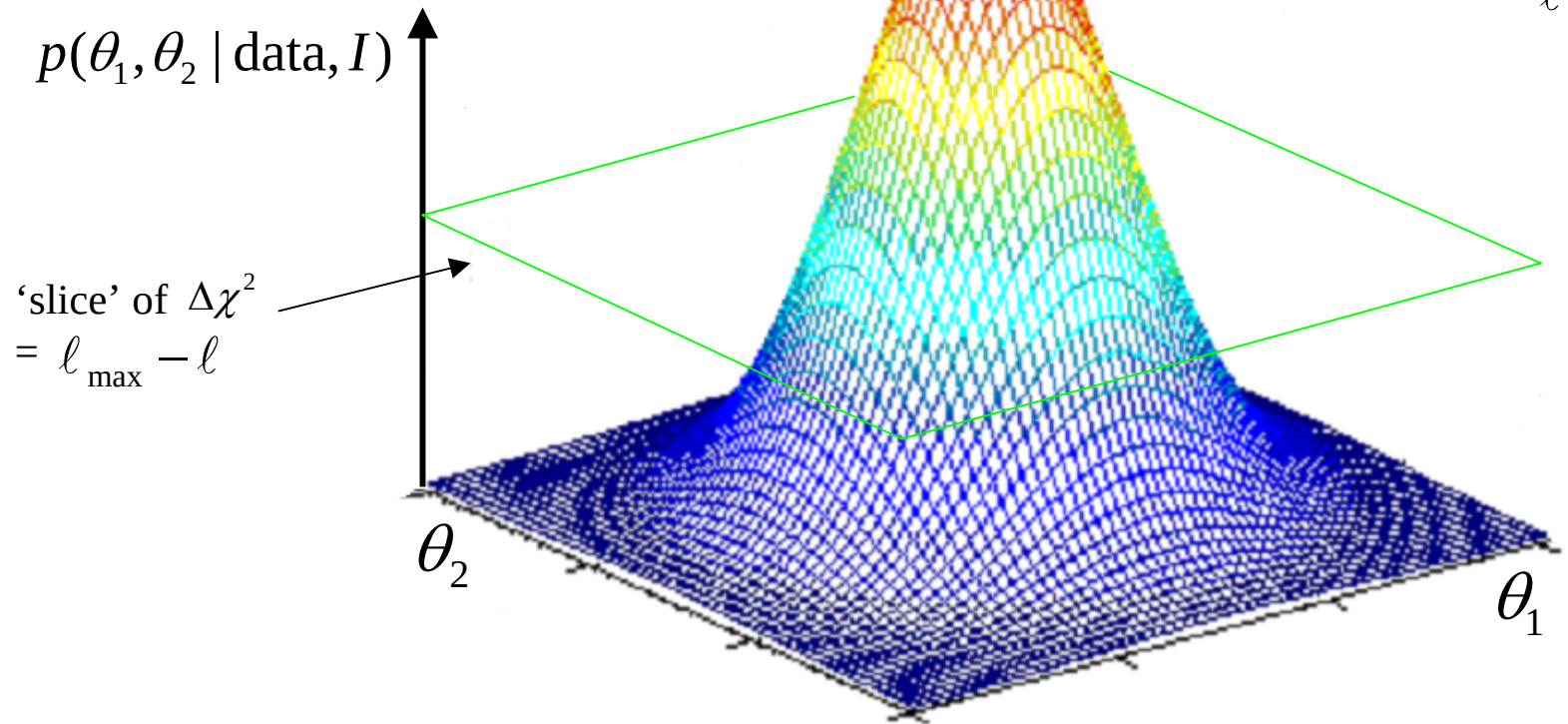
This is a bivariate normal distribution with covariance matrix

$$\sigma_{ij}^2 = \text{cov}_{ij} = \langle (\theta_i - \theta_{i0})(\theta_j - \theta_{j0}) \rangle = \left[-\frac{\partial^2 \ell}{\partial \theta_i \partial \theta_j} \right]^{-1}$$

Fisher information matrix



Parameter estimation: 2-D case



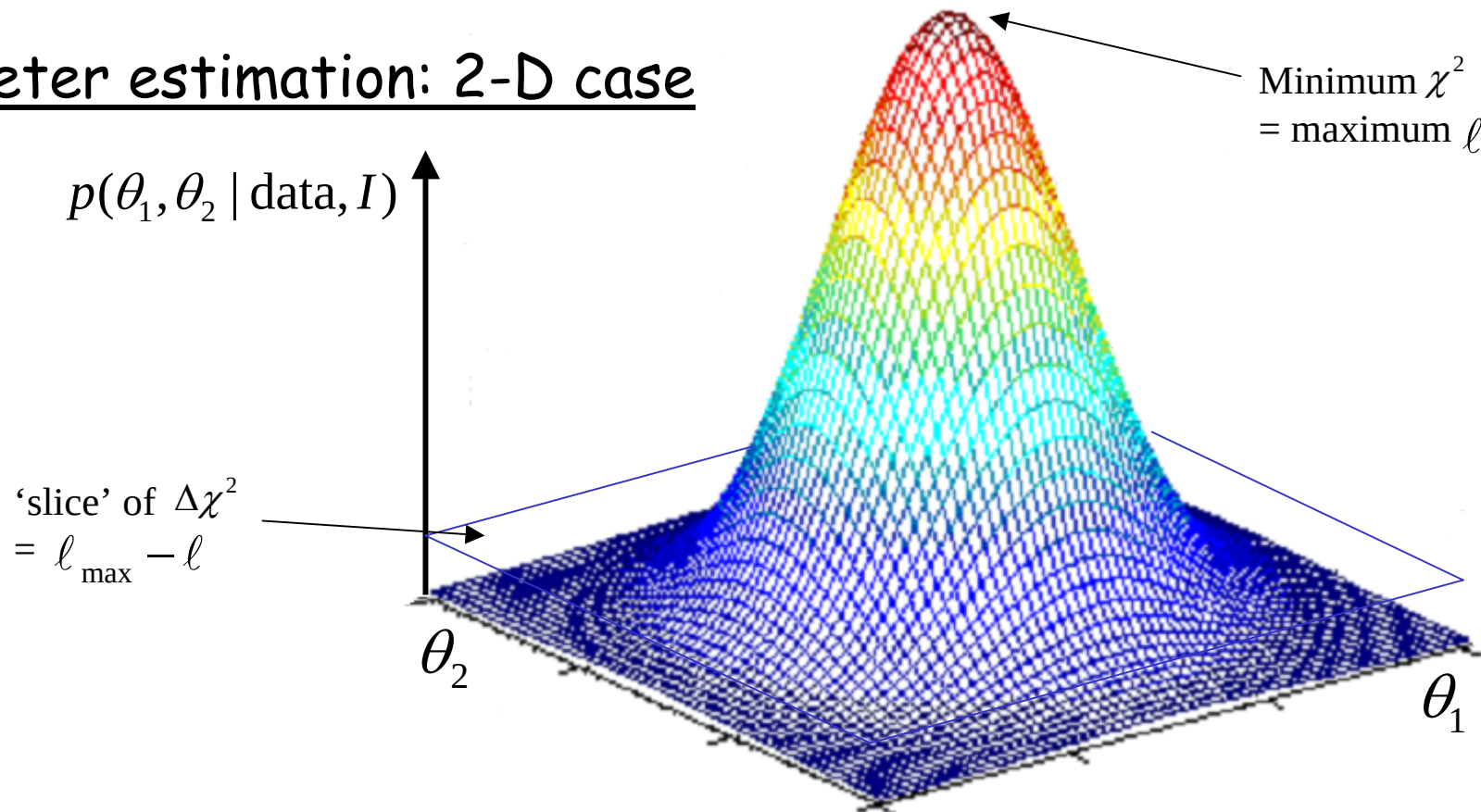
This is a bivariate normal distribution with covariance matrix

$$\sigma_{ij}^2 = \text{cov}_{ij} = \langle (\theta_i - \theta_{i0})(\theta_j - \theta_{j0}) \rangle = \left[-\frac{\partial^2 \ell}{\partial \theta_i \partial \theta_j} \right]^{-1}$$

Fisher information matrix



Parameter estimation: 2-D case



This is a bivariate normal distribution with covariance matrix

$$\sigma_{ij}^2 = \text{cov}_{ij} = \langle (\theta_i - \theta_{i0})(\theta_j - \theta_{j0}) \rangle = \left[-\frac{\partial^2 \ell}{\partial \theta_i \partial \theta_j} \right]^{-1}$$

Fisher information matrix



Parameter estimation: 2-D case

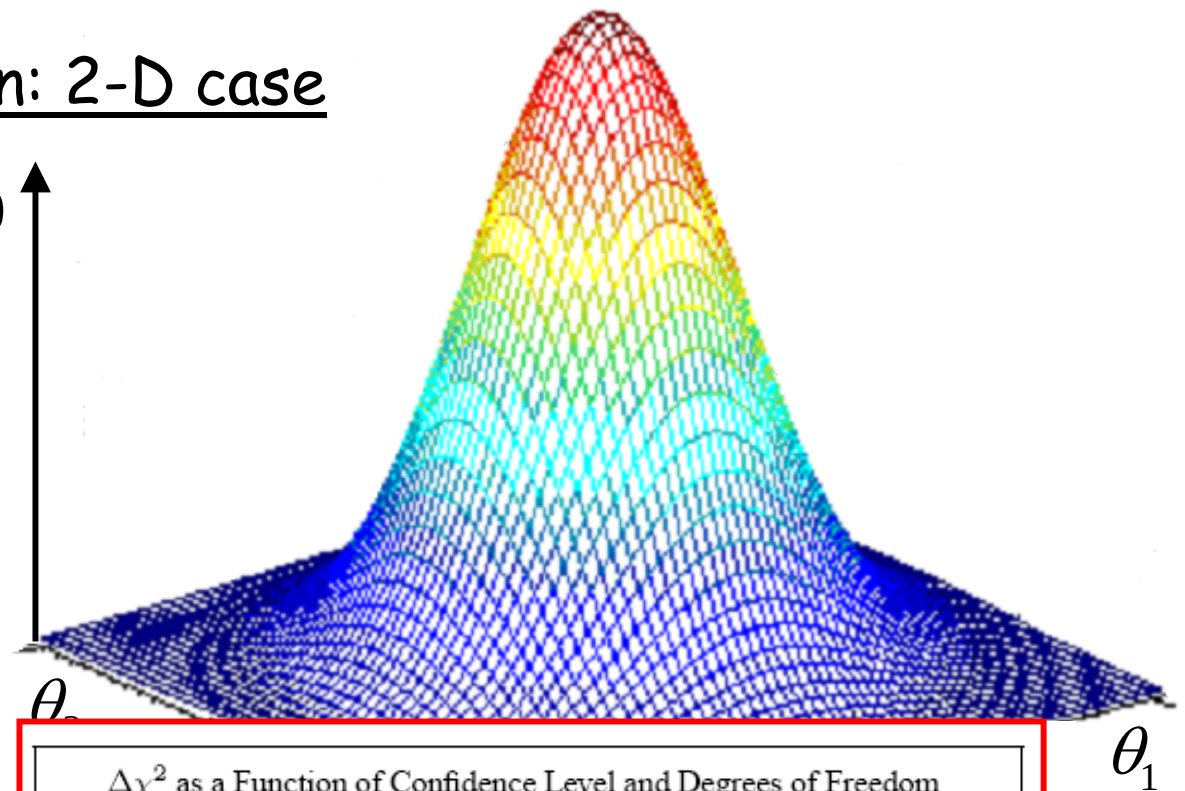
$$p(\theta_1, \theta_2 | \text{data}, I)$$

We can compute the $\Delta\chi^2$ that corresponds to e.g. 68%, 95%, 99% of the posterior pdf.

We can draw contours of equal probability

⇒ **Confidence regions for the parameters**

Extends easily to N parameters - or *degrees of freedom*



$\Delta\chi^2$ as a Function of Confidence Level and Degrees of Freedom						
p	ν					
	1	2	3	4	5	6
68.3%	1.00	2.30	3.53	4.72	5.89	7.04
90%	2.71	4.61	6.25	7.78	9.24	10.6
95.4%	4.00	6.17	8.02	9.70	11.3	12.8
99%	6.63	9.21	11.3	13.3	15.1	16.8
99.73%	9.00	11.8	14.2	16.3	18.2	20.1
99.99%	15.1	18.4	21.1	23.5	25.7	27.8

From Numerical Recipes



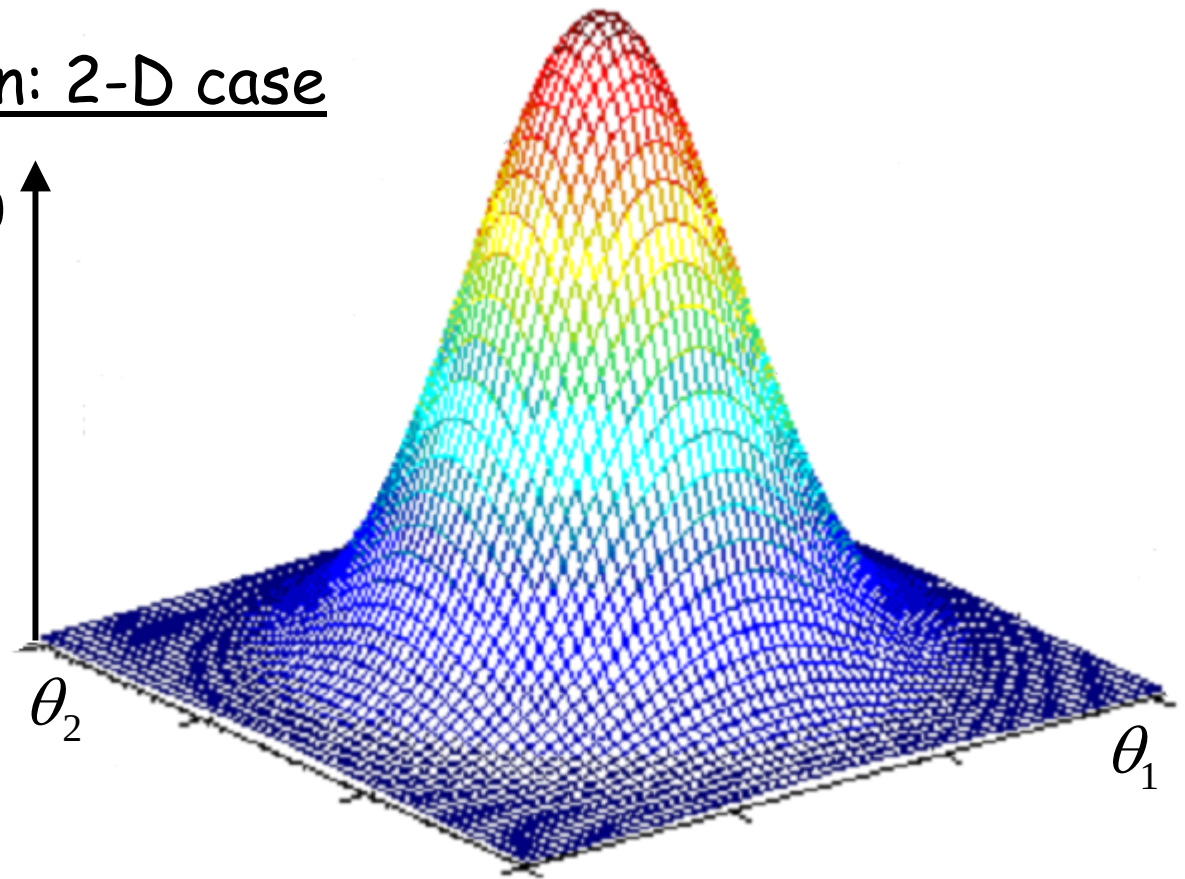
Parameter estimation: 2-D case

$$p(\theta_1, \theta_2 \mid \text{data}, I)$$

Contours of constant probability are **ellipses**.

Covariance matrix is **not** in general diagonal

⇒ What we infer about θ_1 and θ_2 is **not** independent



Parameter estimation: 2-D case

Can define *correlation coefficient*

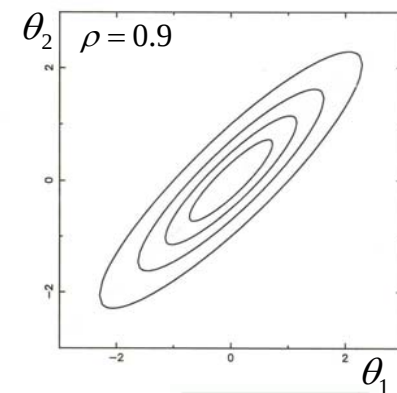
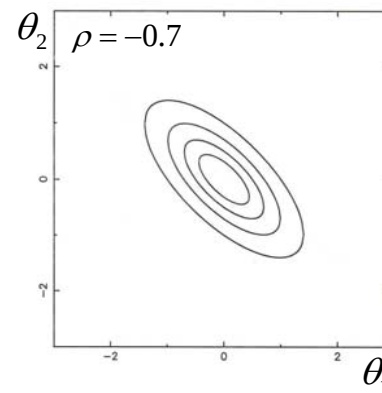
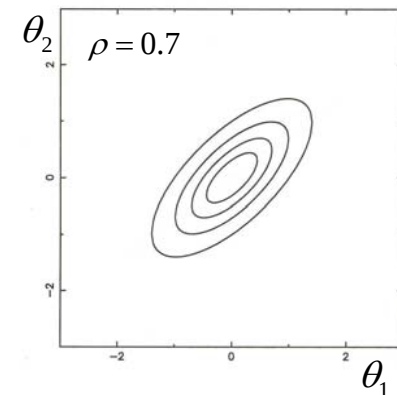
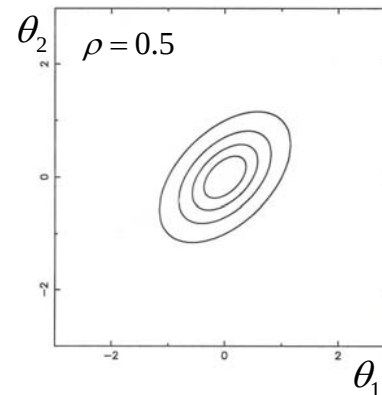
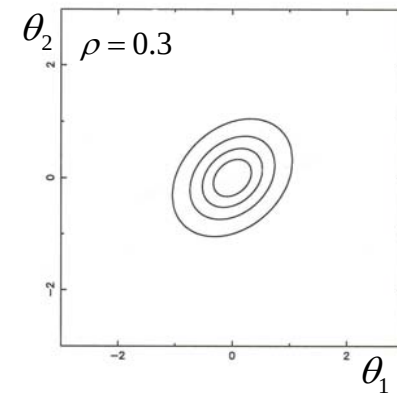
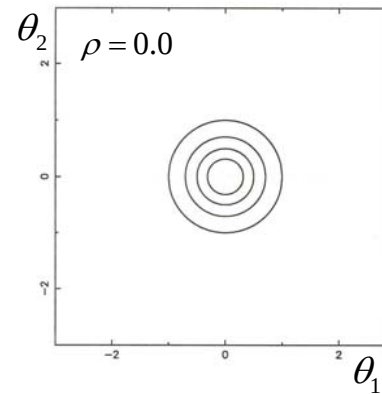
$$\rho = \frac{\text{cov}[\theta_1, \theta_2]}{\sqrt{\text{var}[\theta_1]} \sqrt{\text{var}[\theta_2]}} \quad -1 \leq \rho \leq 1$$

Covariance matrix becomes less diagonal

⇒ $|\rho|$ increases

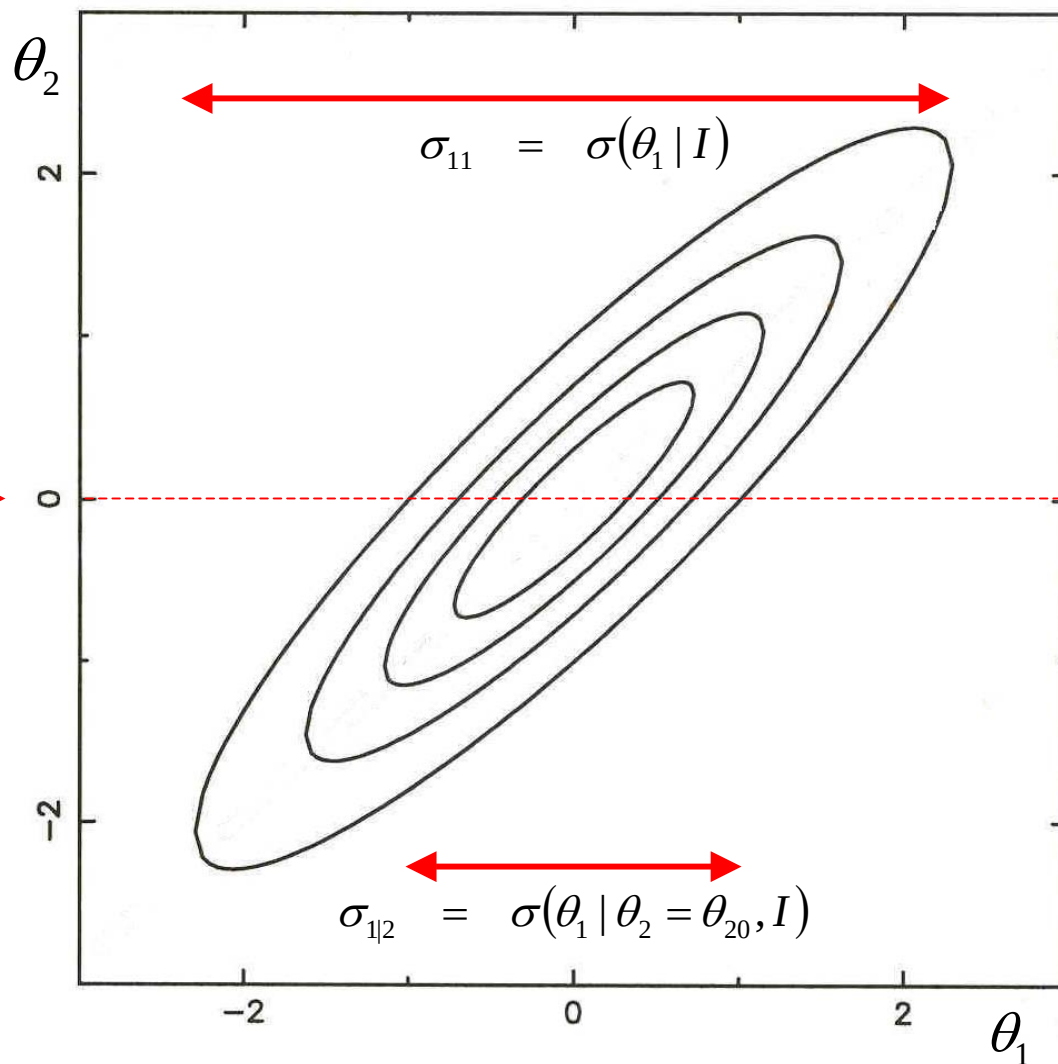
⇒ isoprobability contours elongate

Very important if we are interested only in *one* parameter



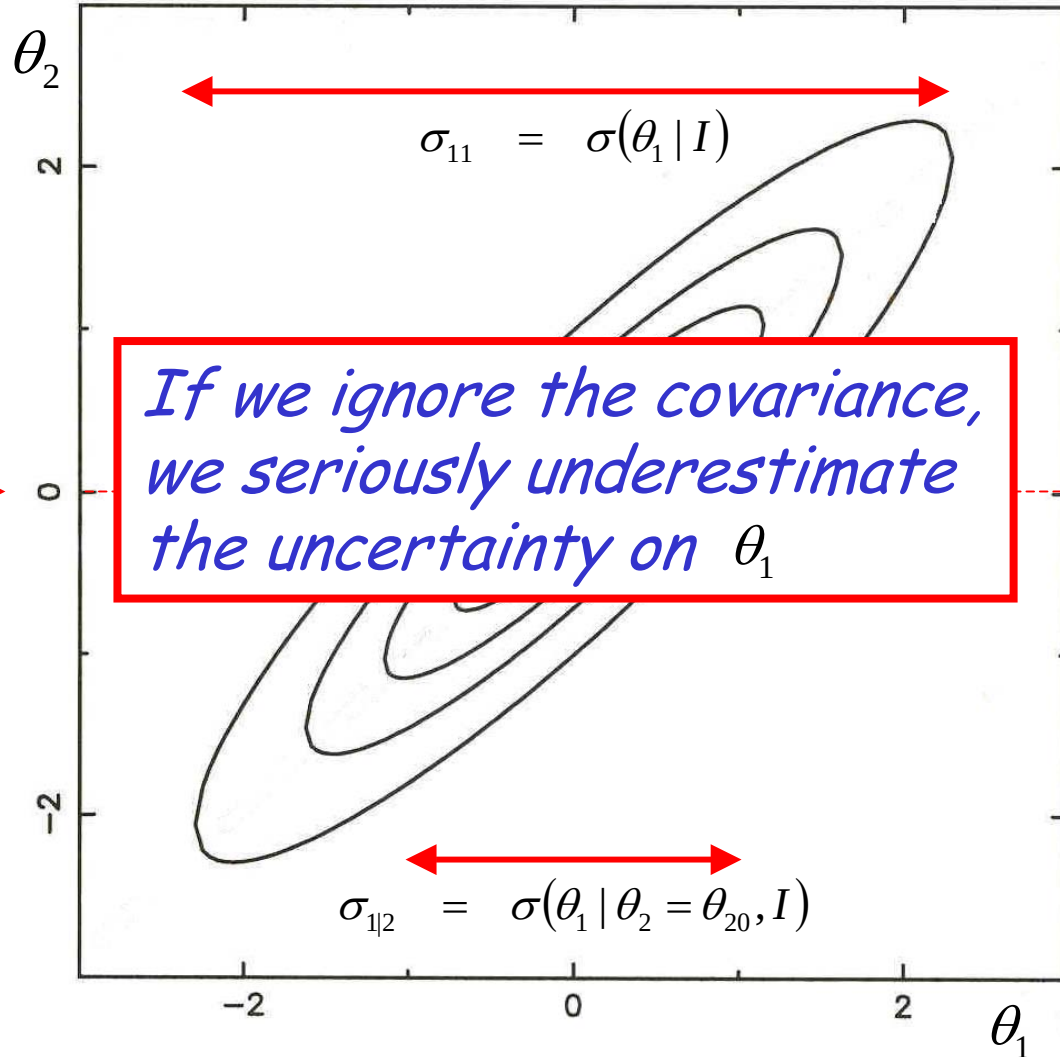
Parameter estimation: 2-D case

'Best-fit' value
of θ_2 , found
from $\left. \frac{\partial \ell}{\partial \theta_j} \right|_{\theta_j = \theta_{0j}} = 0$

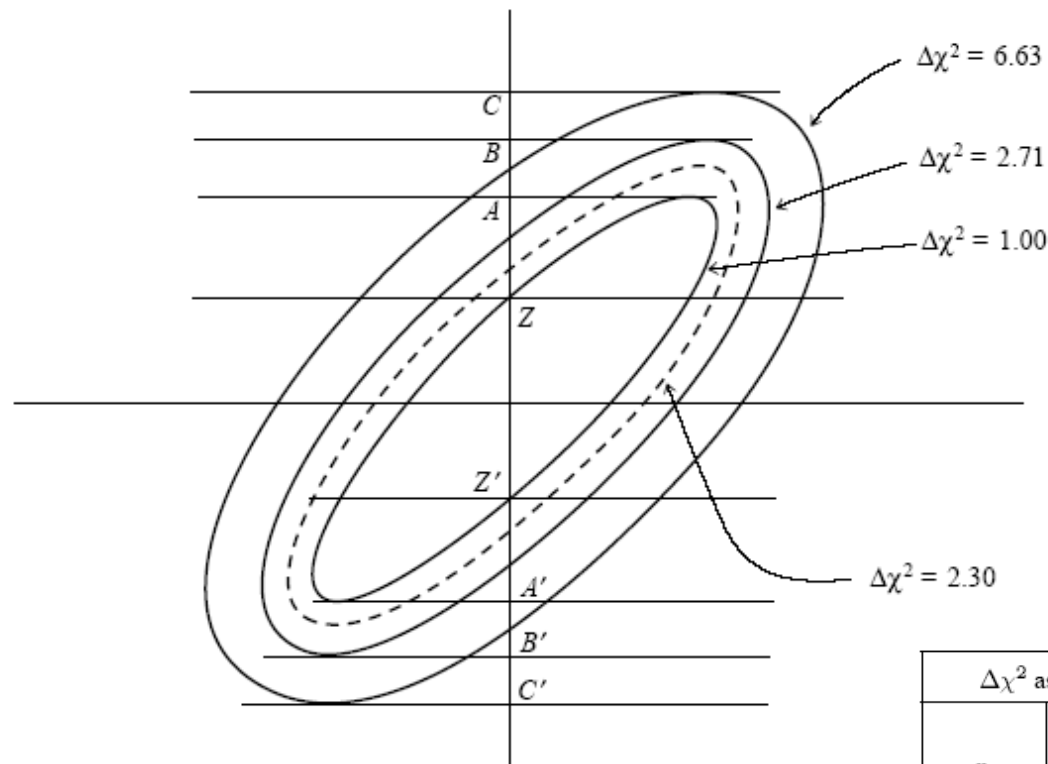


Parameter estimation: 2-D case

'Best-fit' value
of θ_2 , found
from $\left. \frac{\partial \ell}{\partial \theta_j} \right|_{\theta_j = \theta_{0j}} = 0$



Parameter estimation: 2-D case



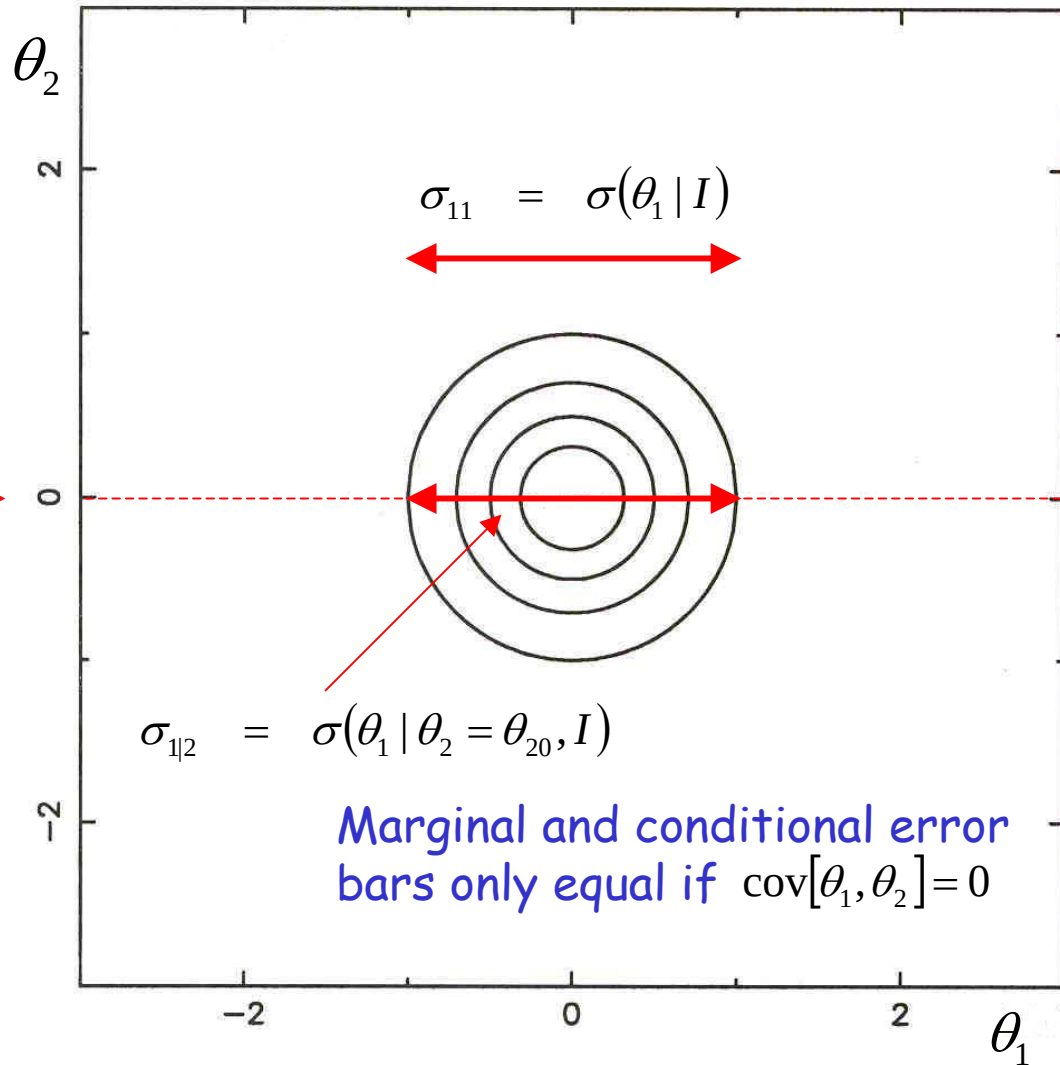
From Numerical Recipes

$\Delta\chi^2$ as a Function of Confidence Level and Degrees of Freedom						
	ν					
p	1	2	3	4	5	6
68.3%	1.00	2.30	3.53	4.72	5.89	7.04
90%	2.71	4.61	6.25	7.78	9.24	10.6
95.4%	4.00	6.17	8.02	9.70	11.3	12.8
99%	6.63	9.21	11.3	13.3	15.1	16.8
99.73%	9.00	11.8	14.2	16.3	18.2	20.1
99.99%	15.1	18.4	21.1	23.5	25.7	27.8



Parameter estimation: 2-D case

'Best-fit' value
of θ_2 , found
from $\left. \frac{\partial \ell}{\partial \theta_j} \right|_{\theta_j = \theta_{0j}} = 0$



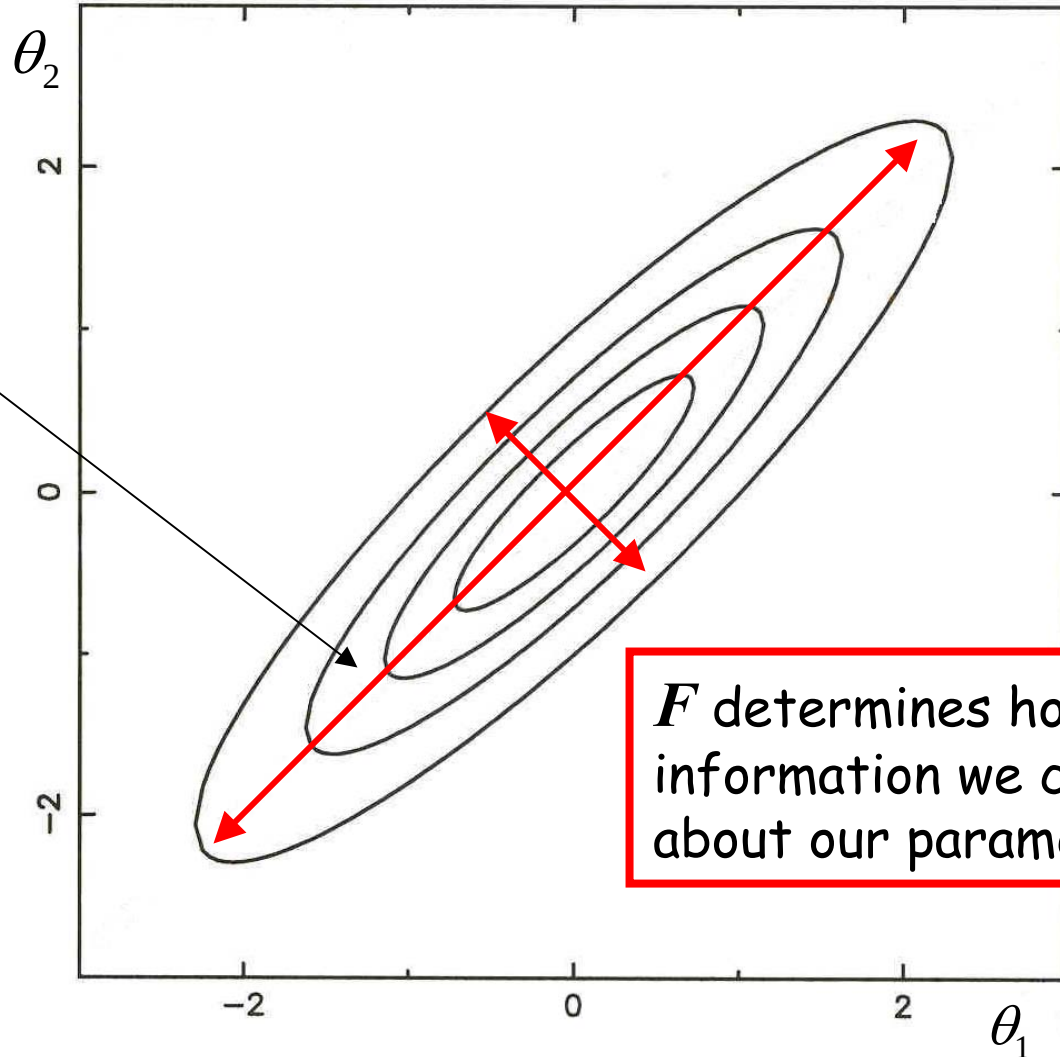
Parameter estimation: 2-D case

Linear combination of θ_1 and θ_2 well constrained by data

Length of axes determined by the **eigenvalues** of the Fisher information matrix

$$F_{ij} = \frac{\partial^2 \ell}{\partial \theta_i \partial \theta_j} = [-\sigma_{ij}^2]^{-1}$$

$$F \theta = \lambda \theta$$



F determines how much information we can learn about our parameters

Direction of axes are the **eigenvectors** of F



Parameter estimation: Gaussian approximation

Taylor expand $\ell(\theta_1, \theta_2)$ around θ_{0j} :

$$\begin{aligned} \ell(\theta_1, \theta_2) = & \ell(\theta_{01}, \theta_{02}) + \frac{\partial \ell}{\partial \theta_1} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01}) + \frac{\partial \ell}{\partial \theta_2} \bigg|_{\theta_j = \theta_{0j}} (\theta_2 - \theta_{02}) + \\ & \frac{1}{2} \left[\frac{\partial^2 \ell}{\partial \theta_1^2} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01})^2 + \frac{\partial^2 \ell}{\partial \theta_2^2} \bigg|_{\theta_j = \theta_{0j}} (\theta_2 - \theta_{02})^2 + 2 \frac{\partial^2 \ell}{\partial \theta_1 \partial \theta_2} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01})(\theta_2 - \theta_{02}) \right] + \dots \end{aligned}$$

$$p(\theta_1, \theta_2 \mid \text{data}, I) = \exp [\ell(\theta_1, \theta_2)]$$

$$= \exp \left[-\frac{1}{2} Q \right] \quad \leftarrow \text{Gaussian approximation}$$

Is the Gaussian approximation a good idea?



Parameter estimation: Gaussian approximation

Taylor expand $\ell(\theta_1, \theta_2)$ around θ_{0j} :

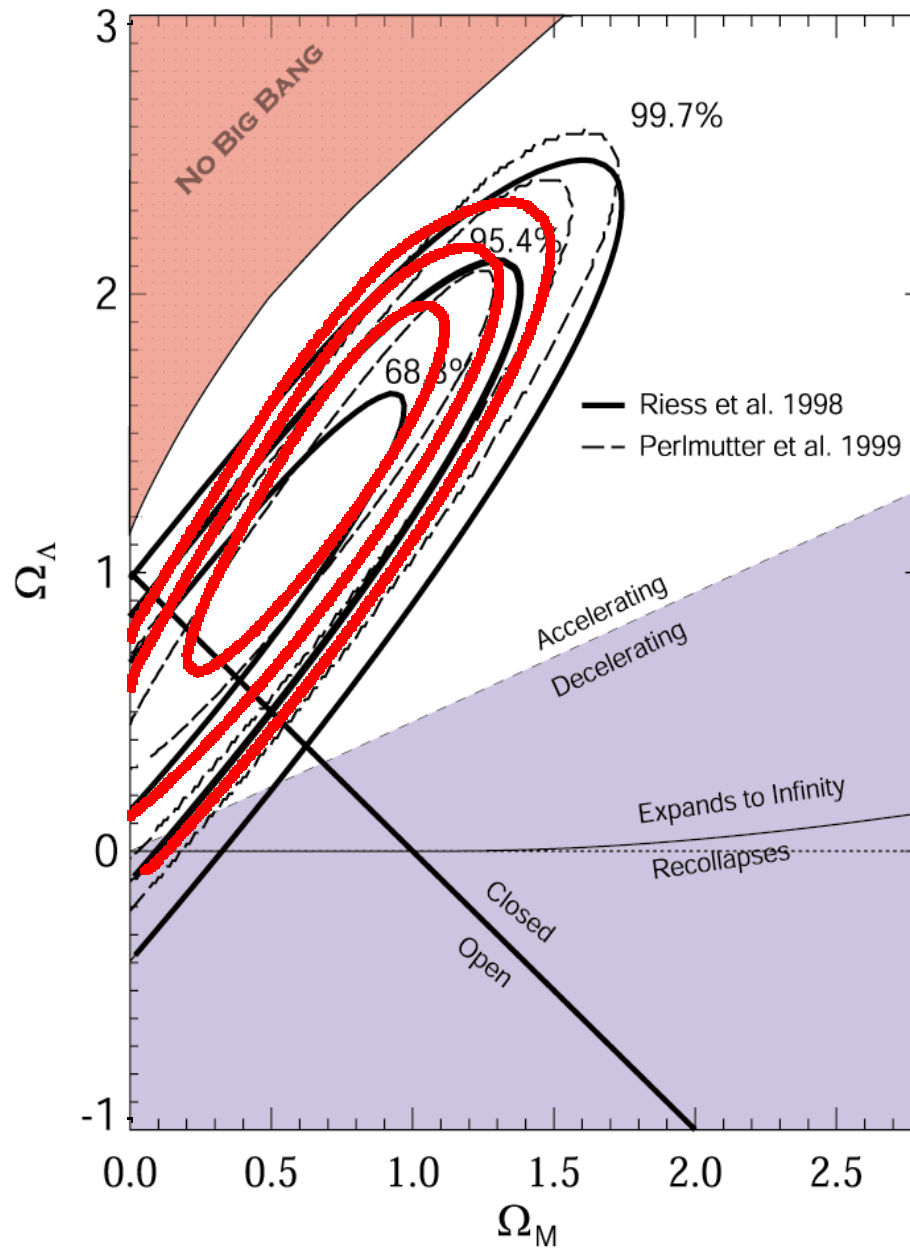
$$\begin{aligned} \ell(\theta_1, \theta_2) = & \ell(\theta_{01}, \theta_{02}) + \frac{\partial \ell}{\partial \theta_1} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01}) + \frac{\partial \ell}{\partial \theta_2} \bigg|_{\theta_j = \theta_{0j}} (\theta_2 - \theta_{02}) + \\ & \frac{1}{2} \left[\frac{\partial^2 \ell}{\partial \theta_1^2} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01})^2 + \frac{\partial^2 \ell}{\partial \theta_2^2} \bigg|_{\theta_j = \theta_{0j}} (\theta_2 - \theta_{02})^2 + 2 \frac{\partial^2 \ell}{\partial \theta_1 \partial \theta_2} \bigg|_{\theta_j = \theta_{0j}} (\theta_1 - \theta_{01})(\theta_2 - \theta_{02}) \right] + \dots \end{aligned}$$

Is the Gaussian approximation a good idea?

- o Greatly simplifies calculations - only need to compute the elements of the Fisher matrix (covariance matrix)
- o Nowadays much better to compute **full posterior** pdf. Not too hard with present-day computers, even for large N

Markov Chain Monte Carlo Methods



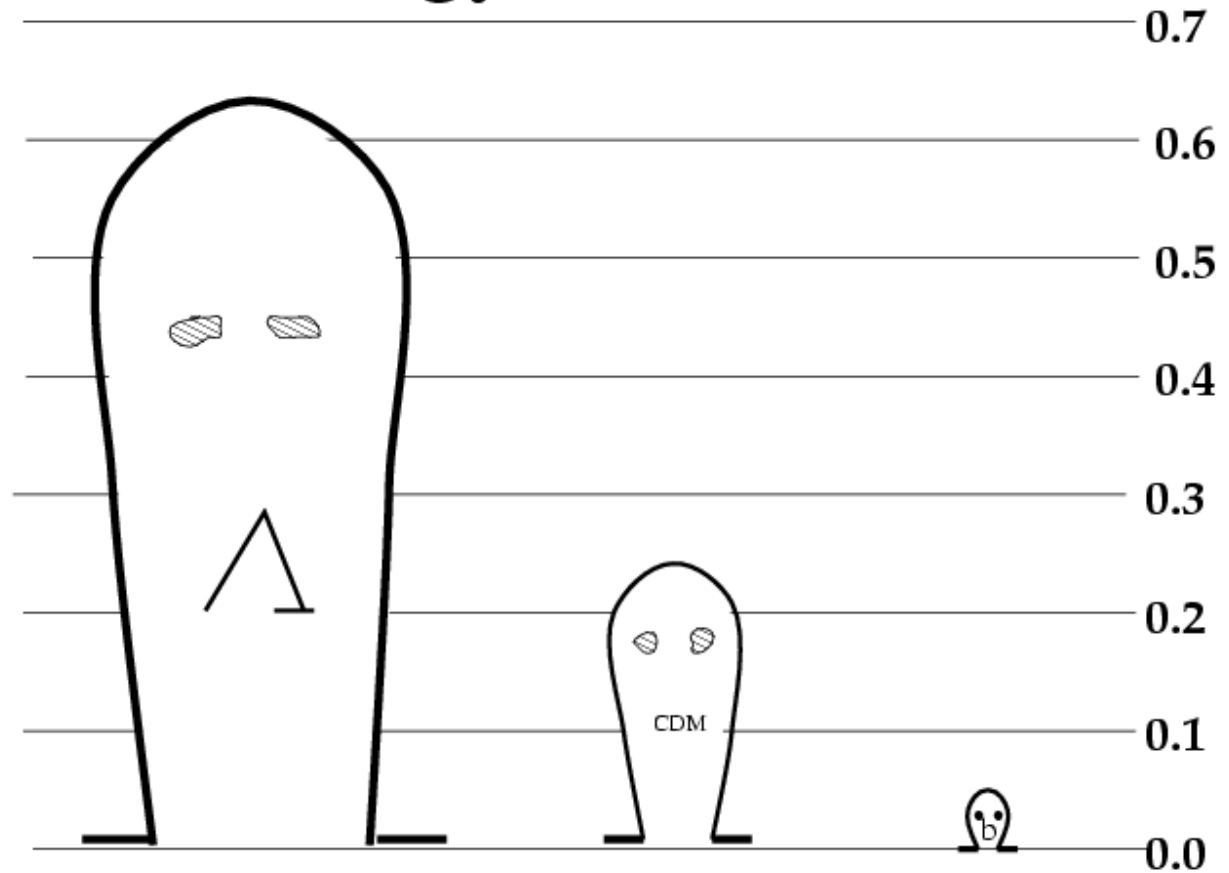


Λ CDM

Cosmology's Most Wanted

Figure 3. A line up of cosmological culprits

Ω_Λ is the big shot controlling the Universe. He's going to make it blow up. Ω_{CDM} would like to make the Universe collapse but can't compete with Ω_Λ . Ω_b just follows Ω_{CDM} around. Like all dangerous criminals, one can never be sure of Ω_Λ until he is behind bars. The CMB police is being beefed up. Hundreds of heroic CMB observers are now planning his capture.



Ω_Λ	Ω_{CDM}	Ω_b
cosmological constant energy of the vacuum He never clumps His evil plan is to blow up the Universe	cold dark matter He likes to clump but has never been detected directly His evil plan is to make the Universe collapse	normal baryonic matter a pawn in the cosmic game who just follows CDM around. He thinks he's a complex life form but is really just a bunch of hydrogen

From Lineweaver (1998)

General Relativity:-

Geometry $\longleftrightarrow$ matter / energy

"Spacetime tells matter how to move and matter tells spacetime how to curve"

Einstein's Field Equations

$$G^{\mu\nu} = R^{\mu\nu} - \frac{1}{2} g^{\mu\nu} R = 8\pi G T^{\mu\nu}$$

Einstein tensor

Ricci tensor

Metric tensor

Curvature scalar

Energy-momentum tensor
of gravitating mass-energy

General Relativity:-

Geometry $\longleftrightarrow$ matter / energy

*"Spacetime tells matter how to move and
matter tells spacetime how to curve"*

Einstein's Field Equations

Given $g^{\mu\nu}$ can compute $R^{\mu\nu}$ and R ;

These are *generated* by $T^{\mu\nu}$



Treat Universe as a *perfect fluid*

$$T^{\mu\nu} = (\rho + P)u^\mu u^\nu + Pg^{\mu\nu}$$

Density

Pressure

Four-velocity

N.B. $c = 1$

Solve to give *Friedmann's Equations*

$$H^2 = \left(\frac{\dot{R}}{R} \right)^2 = \frac{8\pi G\rho}{3} - \frac{k}{R^2}$$

$$\frac{\ddot{R}}{R} = -\frac{4\pi G}{3}(\rho + 3P)$$



Einstein originally sought *static solution* i.e. :-

$$\dot{R} = 0 \text{ for all } t$$

But if $\rho, P \geq 0$ can't have $\ddot{R} = 0$

However, GR actually gives

Covariant derivative

$$G^{\mu\nu}_{;\nu} = T^{\mu\nu}_{;\nu} = 0$$

Can add a constant times $g^{\mu\nu}$ to $G^{\mu\nu}$

Einstein's cosmological constant

$$G^{\mu\nu} = R^{\mu\nu} - \frac{1}{2} g^{\mu\nu} R + g^{\mu\nu} \Lambda$$

ISYA. Ifrane, 2nd - 23rd July 2004



Friedmann's Equations now give:-

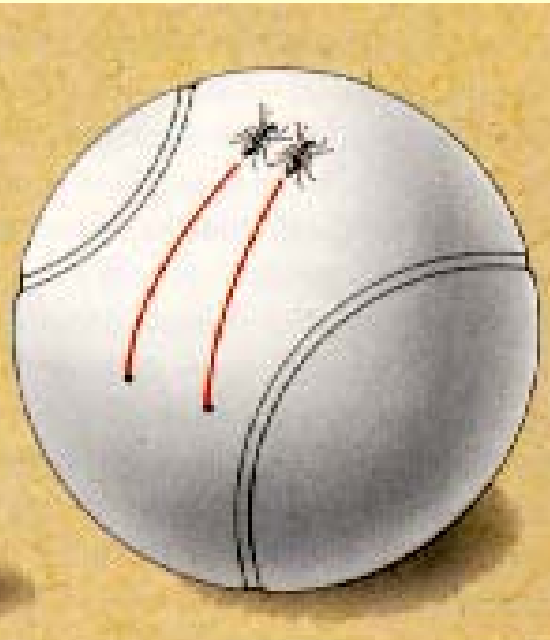
$$H^2 = \left(\frac{\dot{R}}{R} \right)^2 = \frac{8\pi G \rho}{3} + \frac{\Lambda}{3} - \frac{k}{R^2}$$

$$\frac{\ddot{R}}{R} = -\frac{4\pi G}{3}(\rho + 3P) + \frac{\Lambda}{3}$$

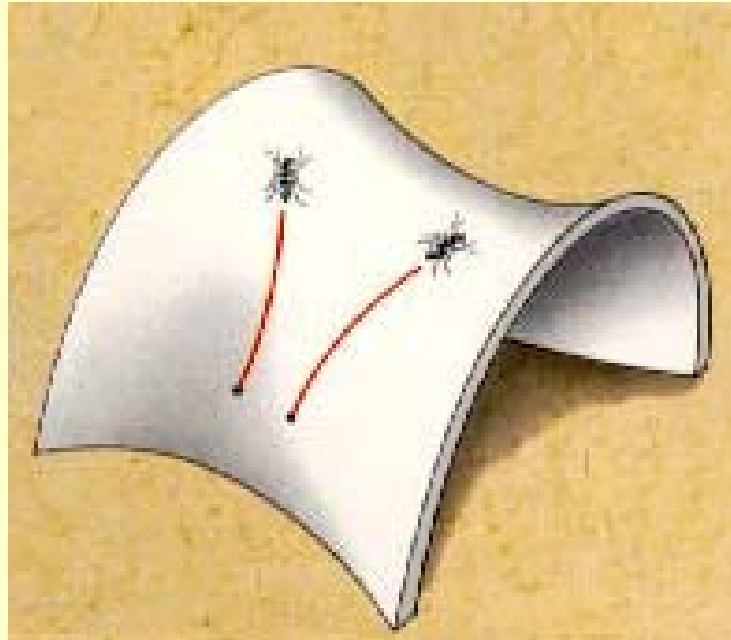
Can tune Λ to give $\dot{R} = 0$ for all t but *unstable*

(and *Hubble expansion* made idea redundant)

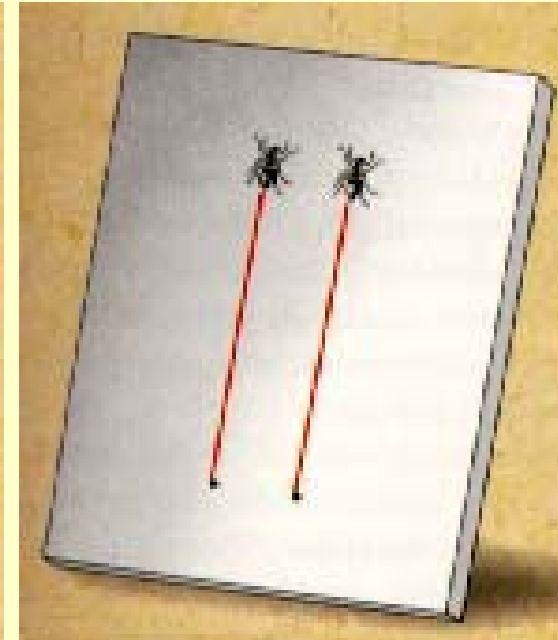




Closed



Open



Flat

$$k = \text{curvature constant} = \begin{cases} -1, & \text{open} \\ 0, & \text{flat} \\ +1, & \text{closed} \end{cases}$$

Einstein's greatest blunder?



Re-expressing Friedmann's Equations:-

For $\Lambda = 0$

$$H^2 = \frac{8\pi G\rho}{3} - \frac{k}{R^2} \quad \Rightarrow \quad k = 0 \Leftrightarrow \rho = \left[\frac{8\pi G}{3H^2} \right]^{-1} = \rho_{\text{crit}}$$

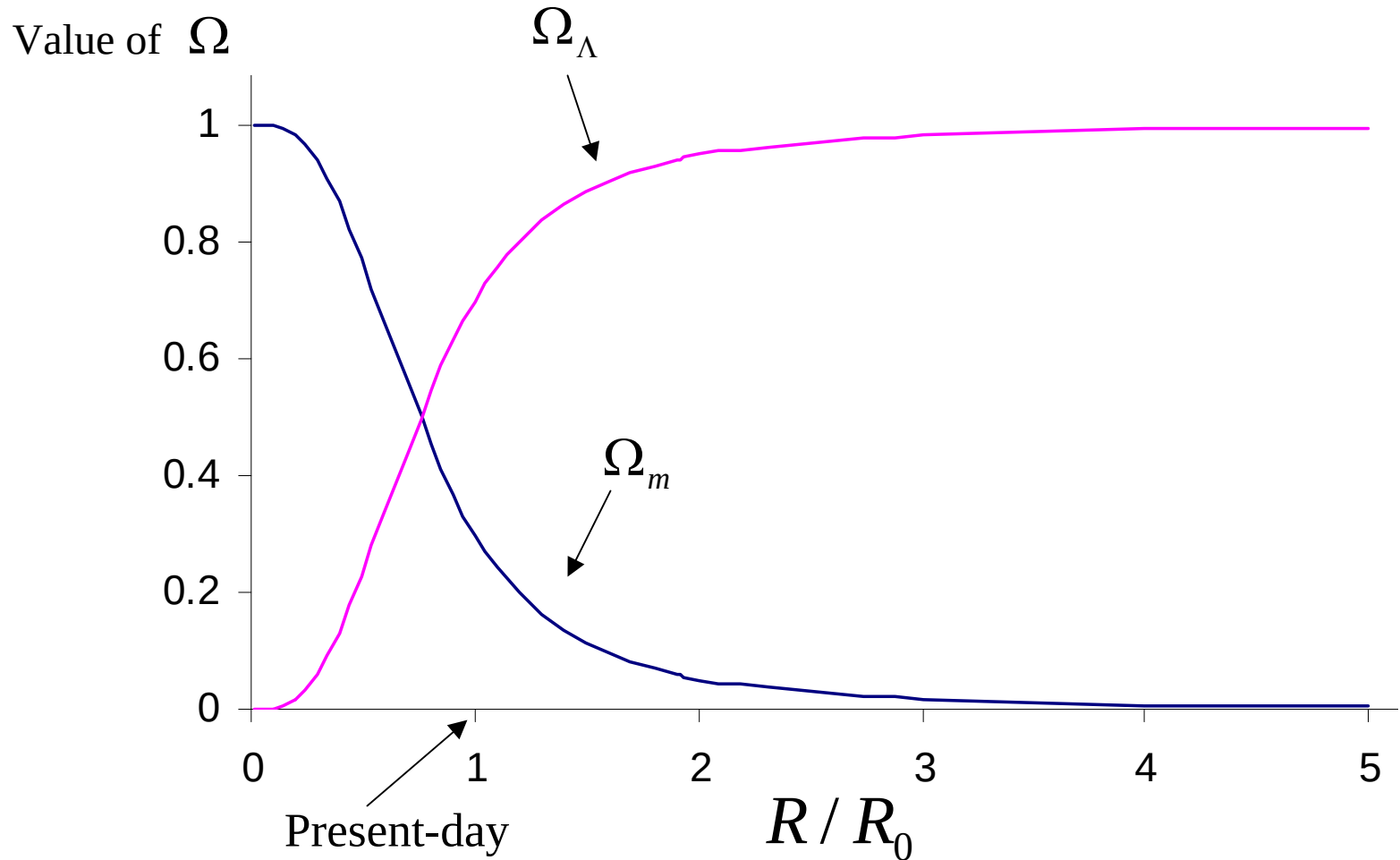
Define

$$\Omega_m = \frac{\rho}{\rho_{\text{crit}}} = \frac{8\pi G\rho}{3H^2} \quad \Omega_\Lambda = \frac{\Lambda}{3H^2} \quad \Omega_k = -\frac{k}{R^2 H^2}$$

It follows that, at *any time*

$$\Omega_m + \Omega_\Lambda + \Omega_k = 1$$





If the Concordance Model is right, we live at a special epoch. Why?...

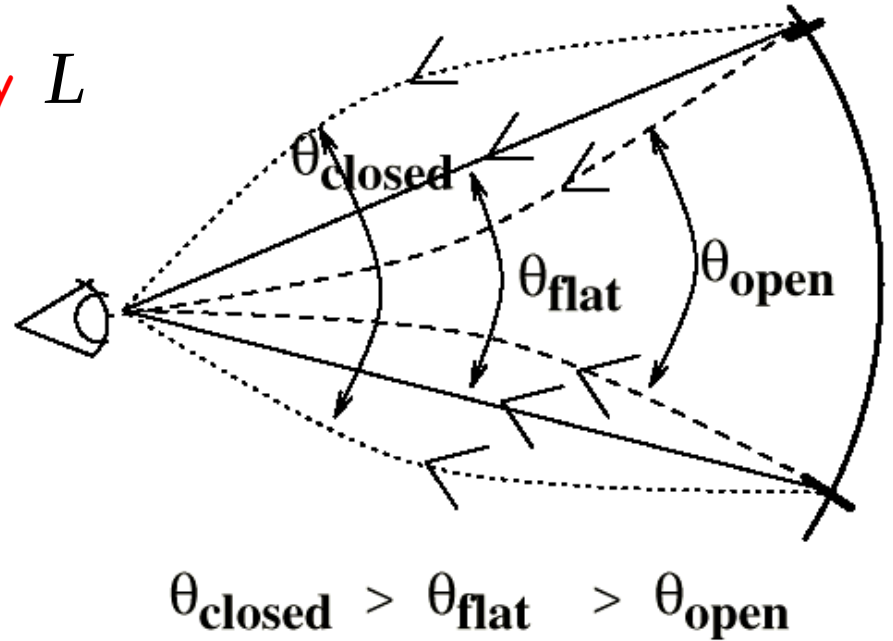
Hubble diagram of distant supernovae

Consider an object of intrinsic luminosity L
from which we observe a flux $\mathfrak{I}$

Define the Luminosity Distance via:-

$$\mathfrak{I} = \frac{L}{4\pi d_L^2}$$

Distance required to give
observed flux *if* Universe
has a flat geometry



Actual distance depends on true
geometry, and expansion history
of the Universe

$$d_L \equiv d_L(z; \Omega_m, \Omega_\Lambda) = (1+z)^2 d_{\text{ang}}(z; \Omega_m, \Omega_\Lambda)$$

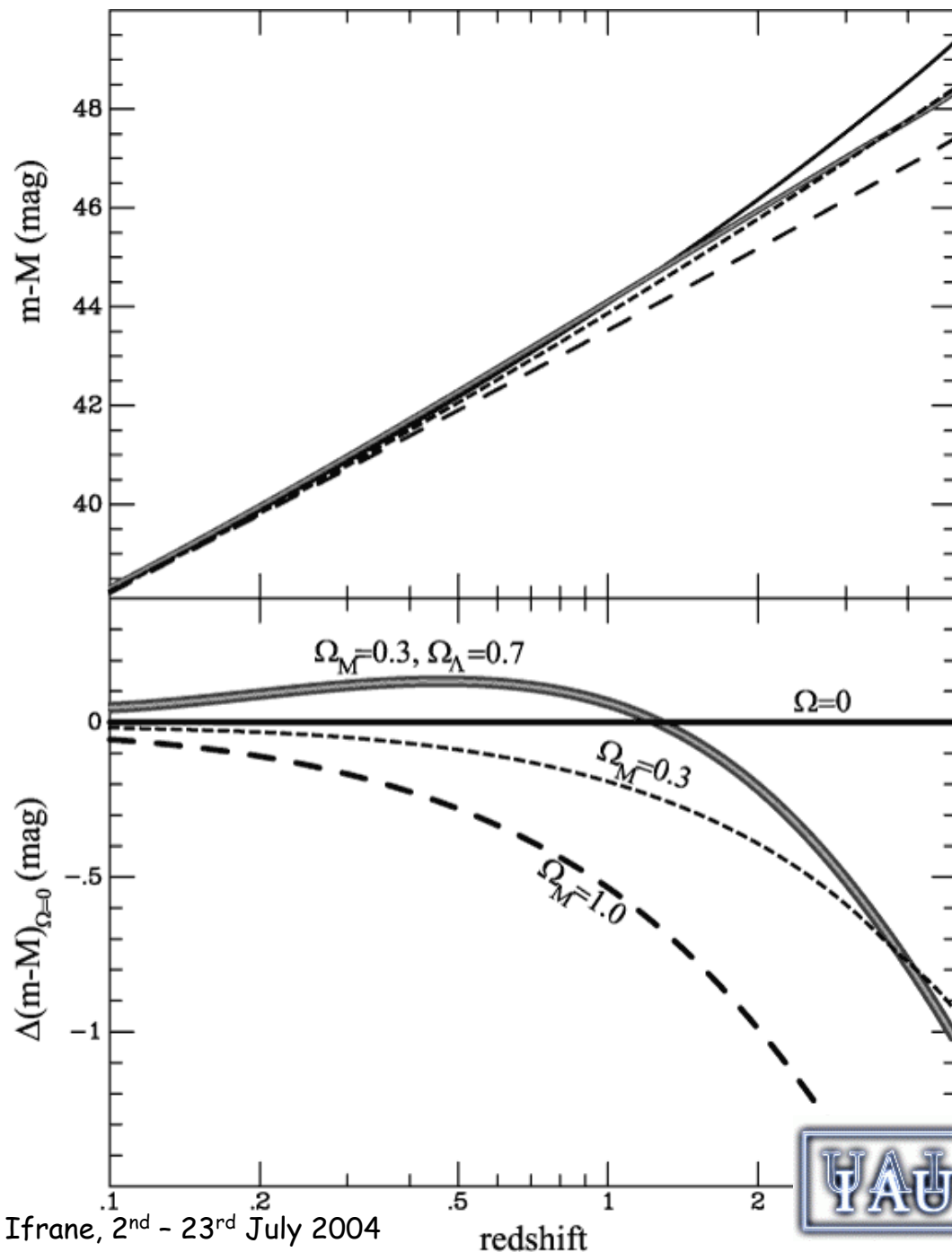
Distance Modulus

$$(m - M)_{mag} = 5 \log \left[\frac{d_L}{Mpc} \right] + 25$$

Fractional distance change
 $\cong \frac{1}{2}(\text{mag change})$

e.g.

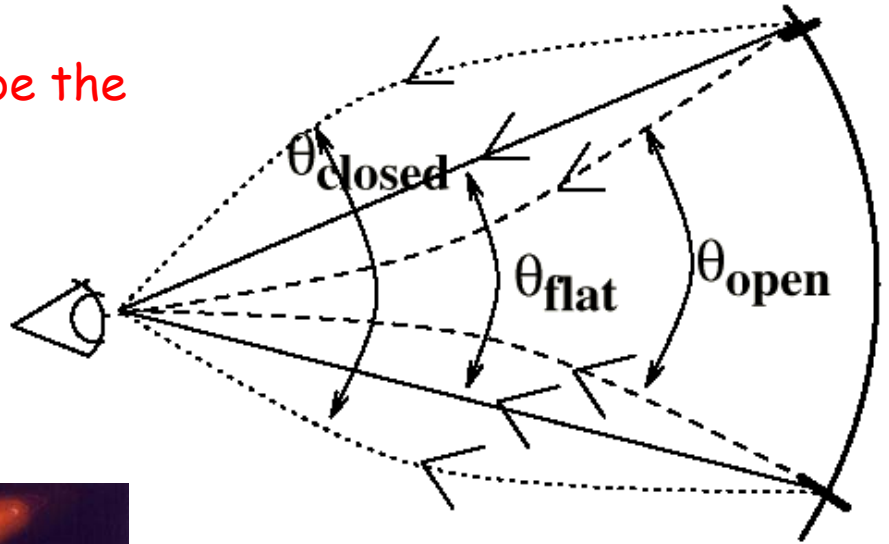
0.1 mag difference is 5%
 distance difference



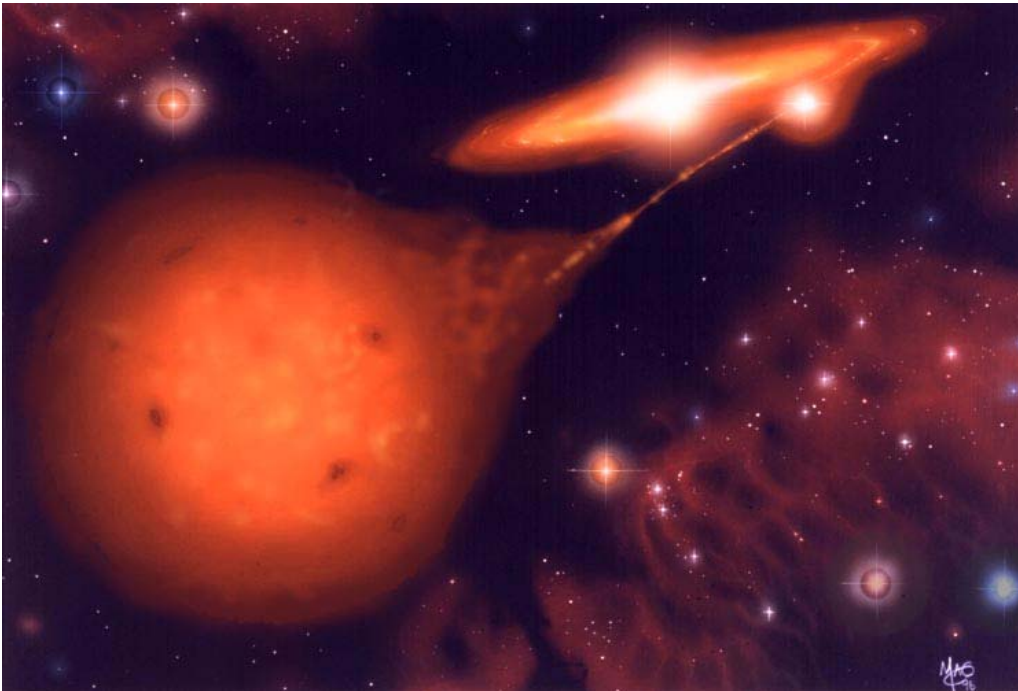
Hubble diagram of distant supernovae

We need a good standard candle, to probe the geometry of the Universe

Type Ia Supernovae

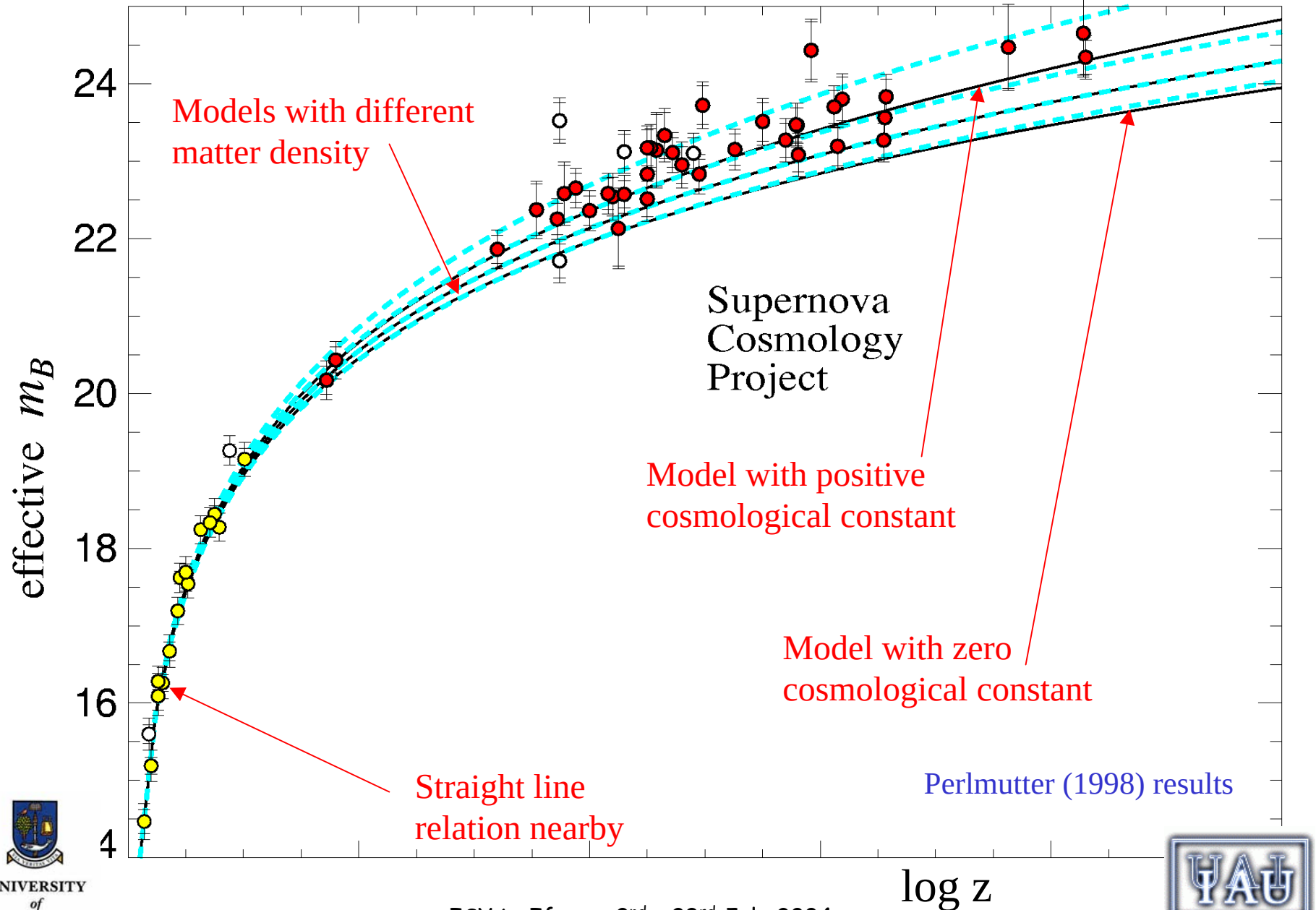


$$\theta_{\text{closed}} > \theta_{\text{flat}} > \theta_{\text{open}}$$



- o Visible to huge distances
- o Small scatter in luminosity at peak brightness
- o Observational programs to detect dozens (100s)
- o Can also use 'Light Curve Shape' to improve distance estimates

Hubble diagram of distant Type Ia supernovae

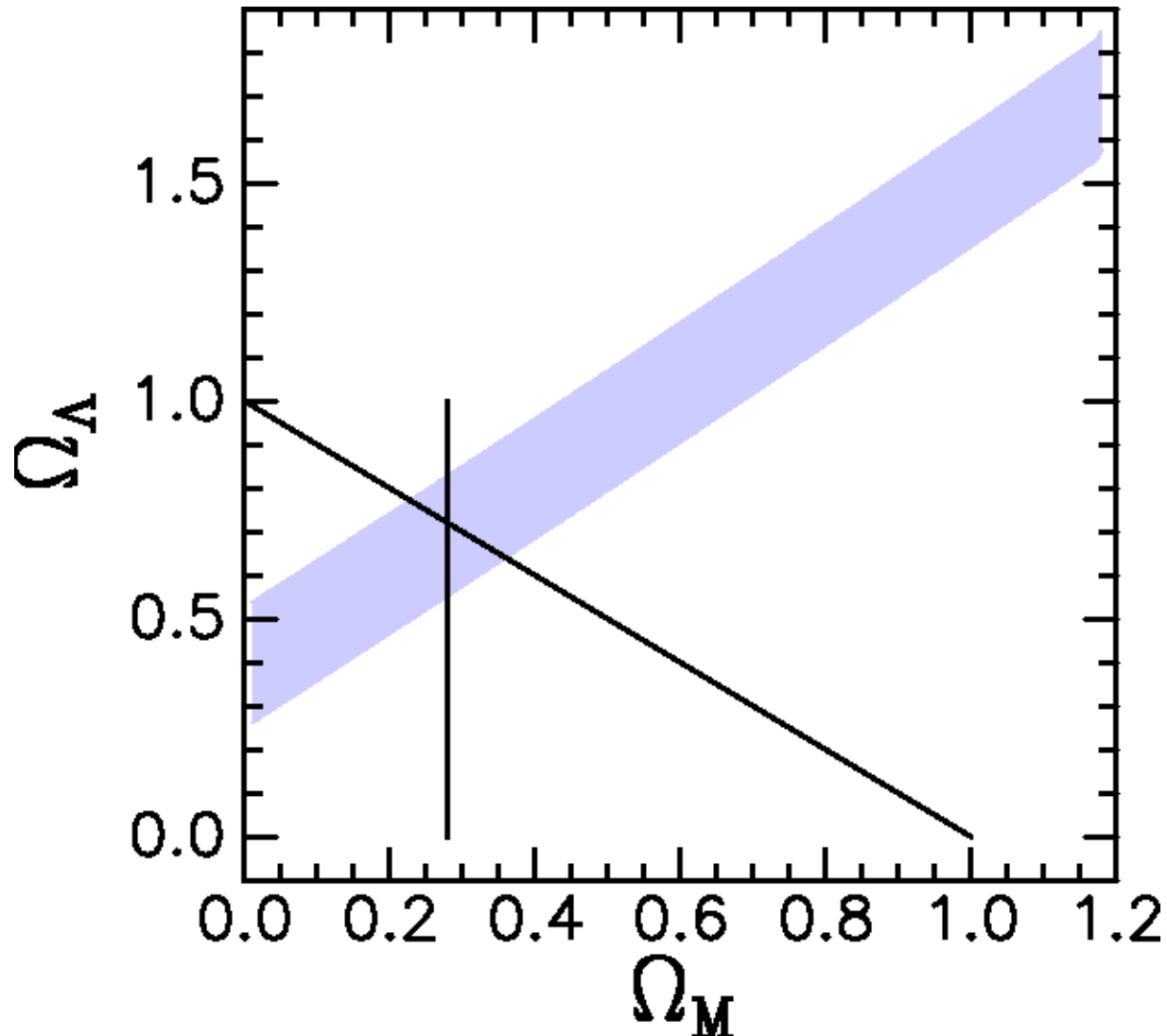


SN Ia at $z = 0.5$

At low redshift, SN Ia essentially measure the **deceleration parameter**

$$q_0 \equiv \frac{-R(t_0)\ddot{R}(t_0)}{\dot{R}^2(t_0)}$$
$$= \frac{1}{2} \sum_i \Omega_i (1 + 3w_i)$$

e.g., $q_0 = \frac{1}{2} \Omega_m - \Omega_\Lambda$

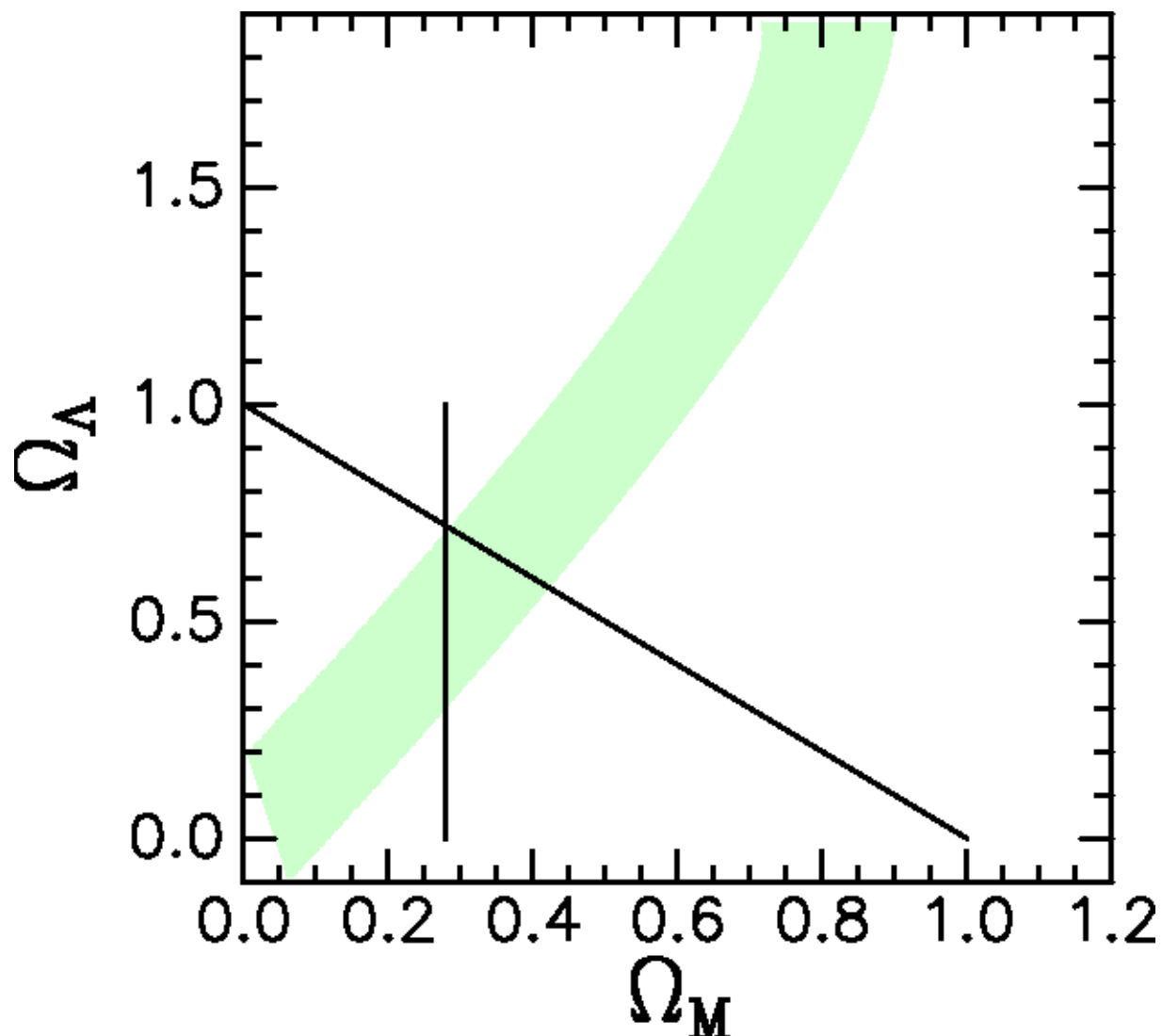


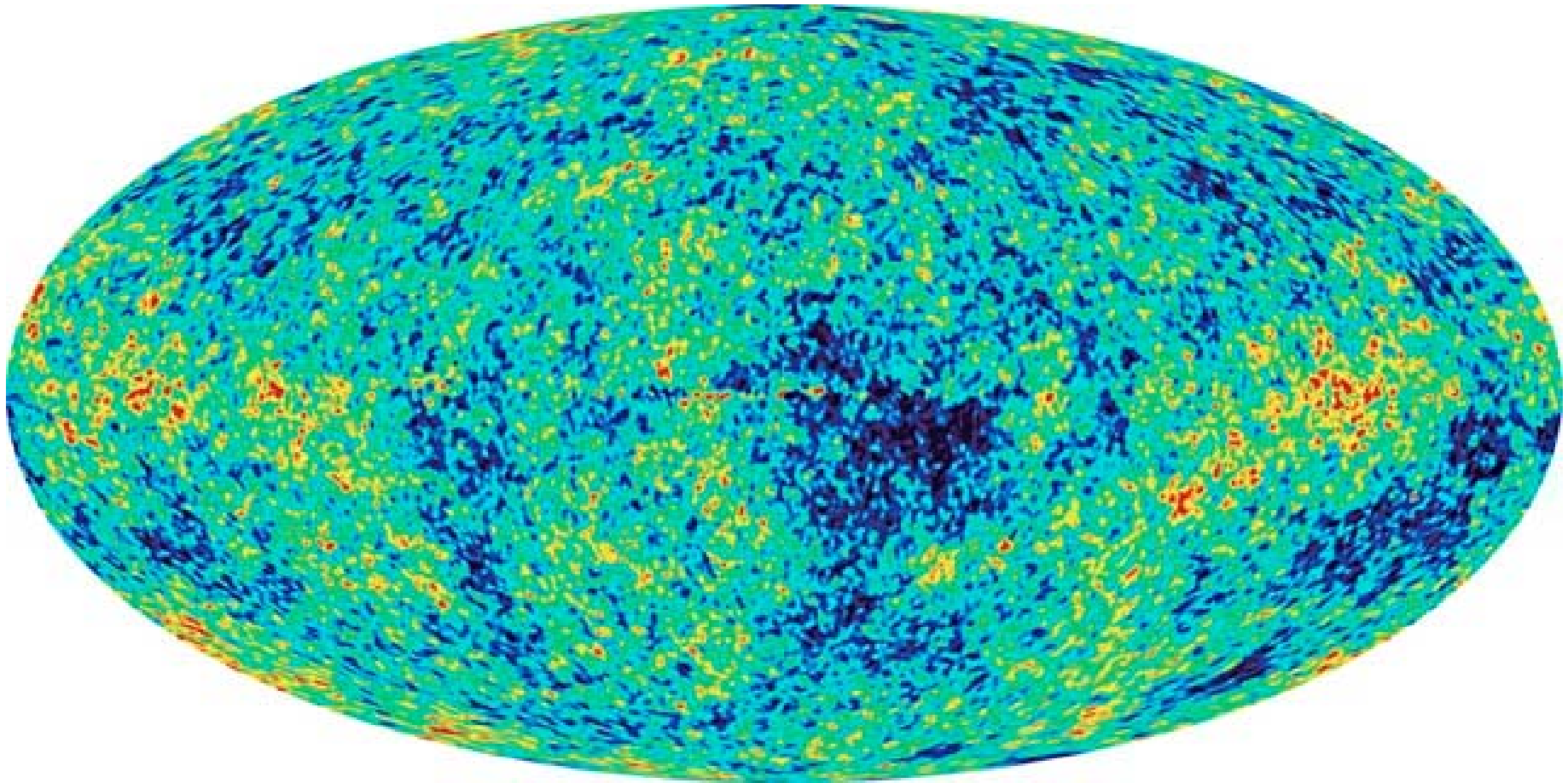
SN Ia at $z = 1.0$

At low redshift, SN Ia essentially measure the **deceleration parameter**

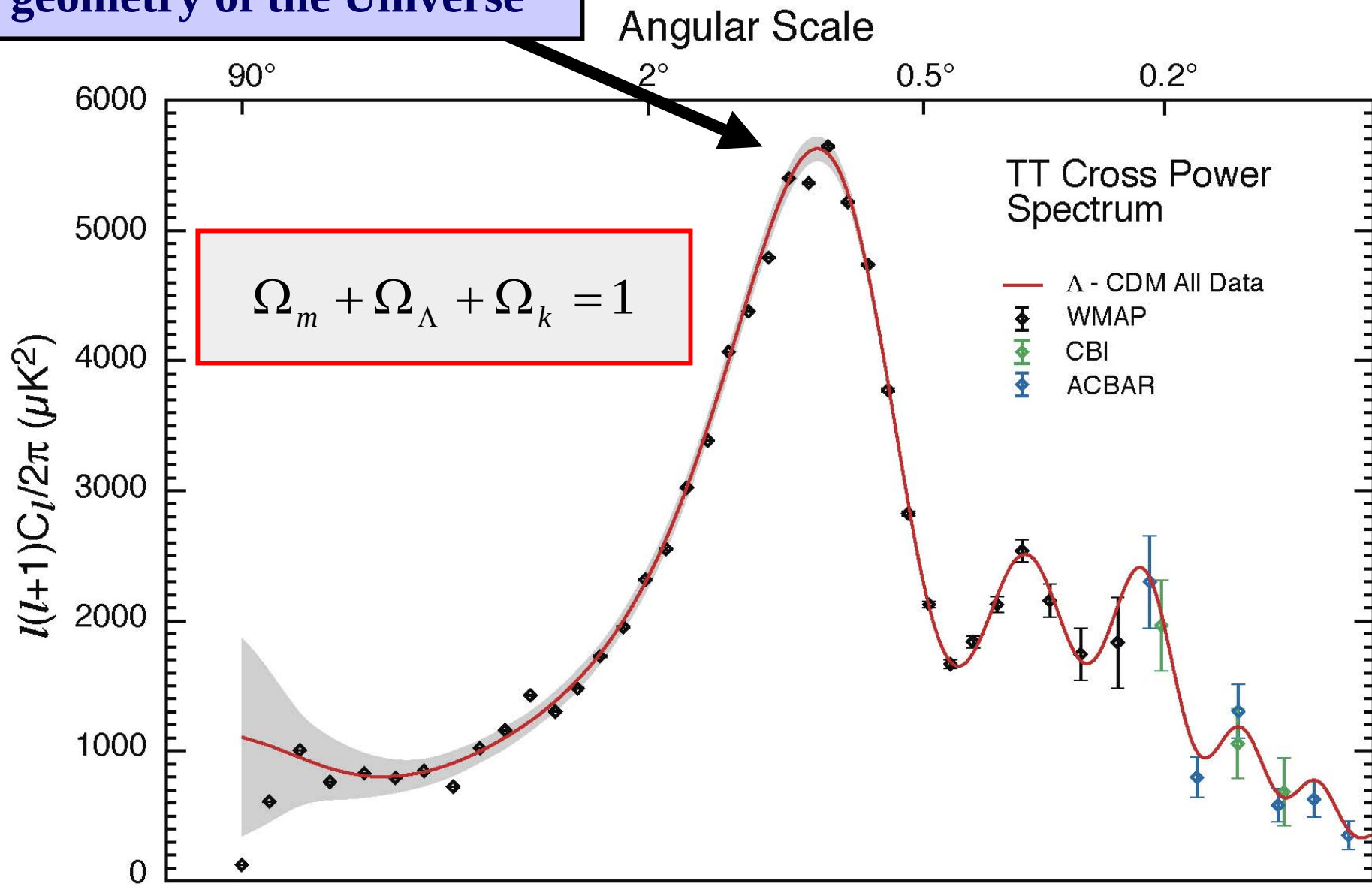
$$q_0 \equiv \frac{-R(t_0)\ddot{R}(t_0)}{\dot{R}^2(t_0)}$$
$$= \frac{1}{2} \sum_i \Omega_i (1 + 3w_i)$$

e.g., $q_0 = \frac{1}{2} \Omega_m - \Omega_\Lambda$





**Position of first peak
sensitive probe of the
geometry of the Universe**



SN Ia measure:-

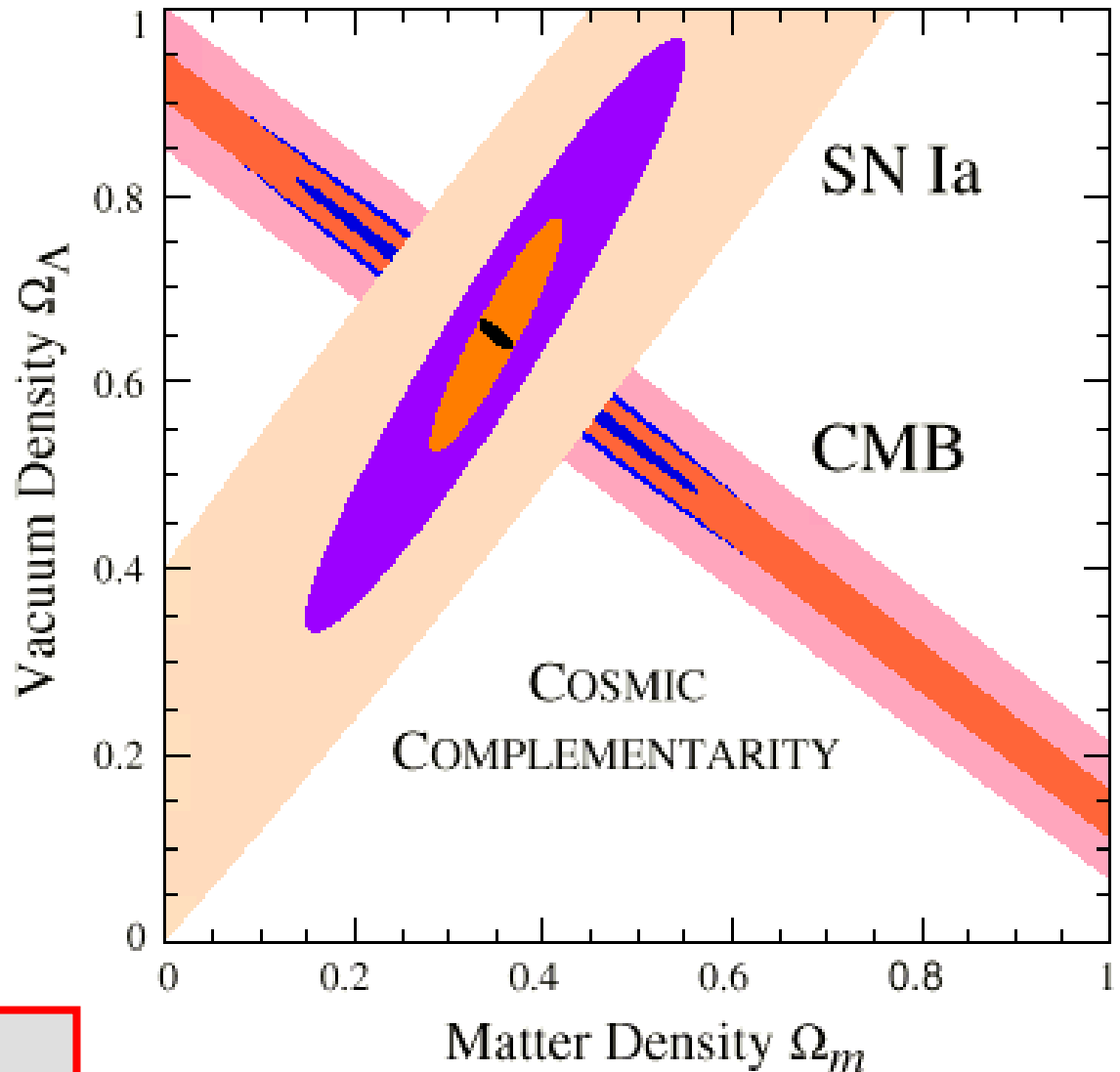
$$q_0 = \frac{1}{2}\Omega_m - \Omega_\Lambda$$

CMBR measures:-

$$\Omega_k = 1 - \Omega_m - \Omega_\Lambda$$

Together, can
constrain:-

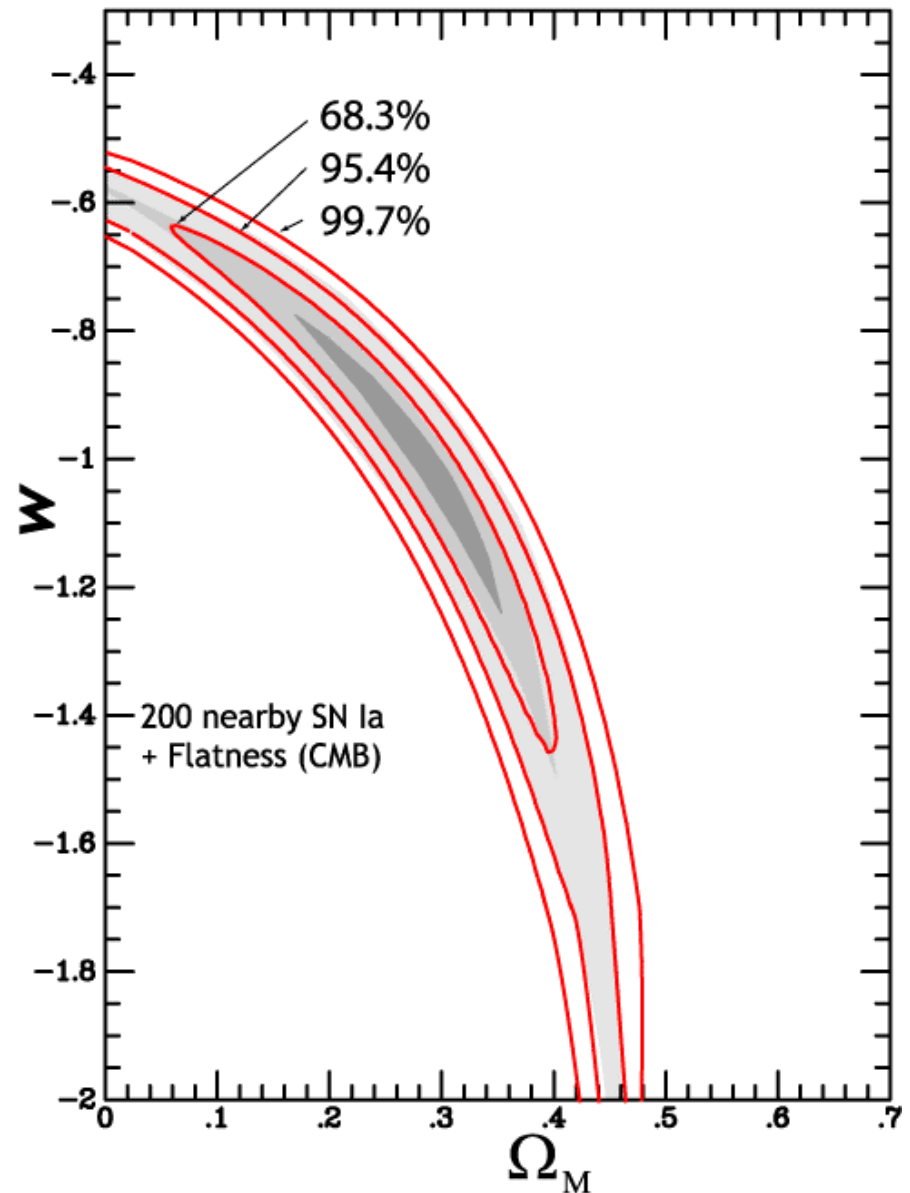
$$\Omega_m, \Omega_\Lambda$$



Can we distinguish a constant Λ term from quintessence?

Not from current ground-based SN observations (combined with e.g. LSS)...

...or from future ground-based observations



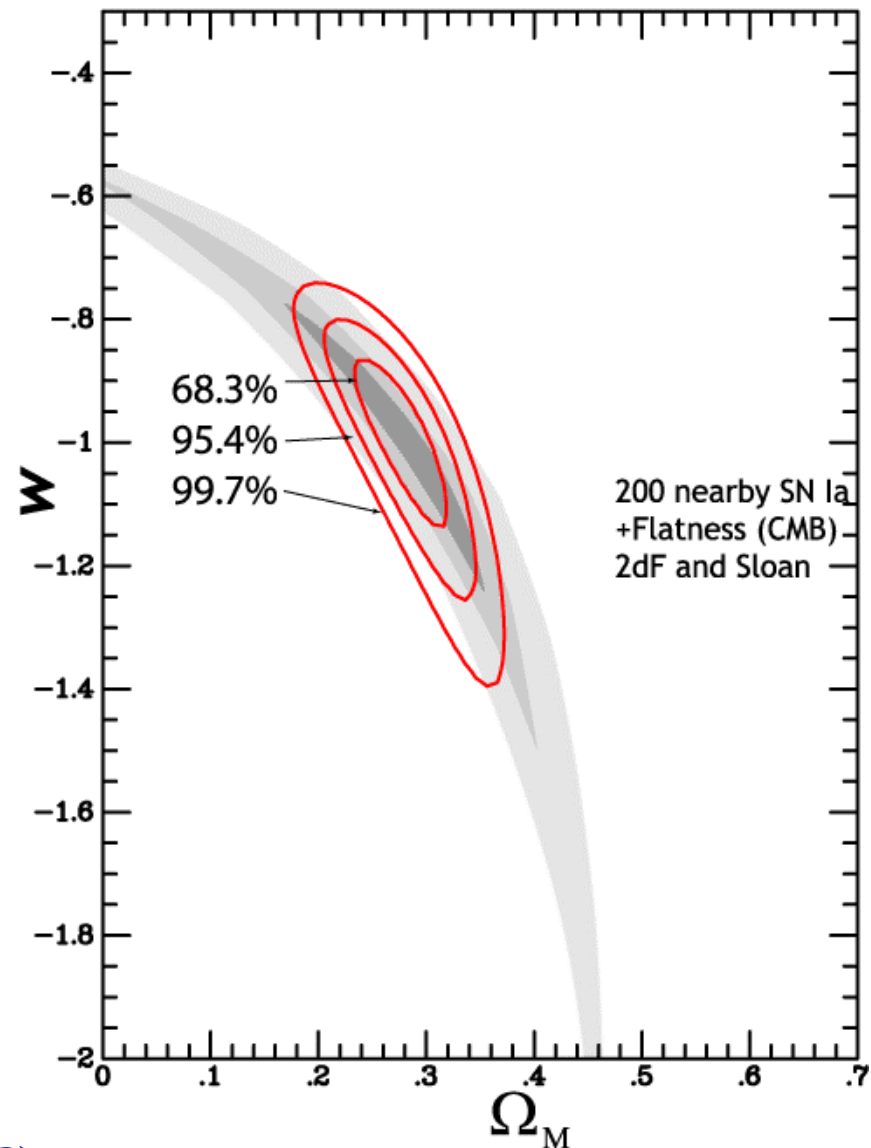
Adapted from Schmidt (2002)



Can we distinguish a constant Λ term from quintessence?

Not from current ground-based SN observations (combined with e.g. LSS)...

...or from future ground-based observations



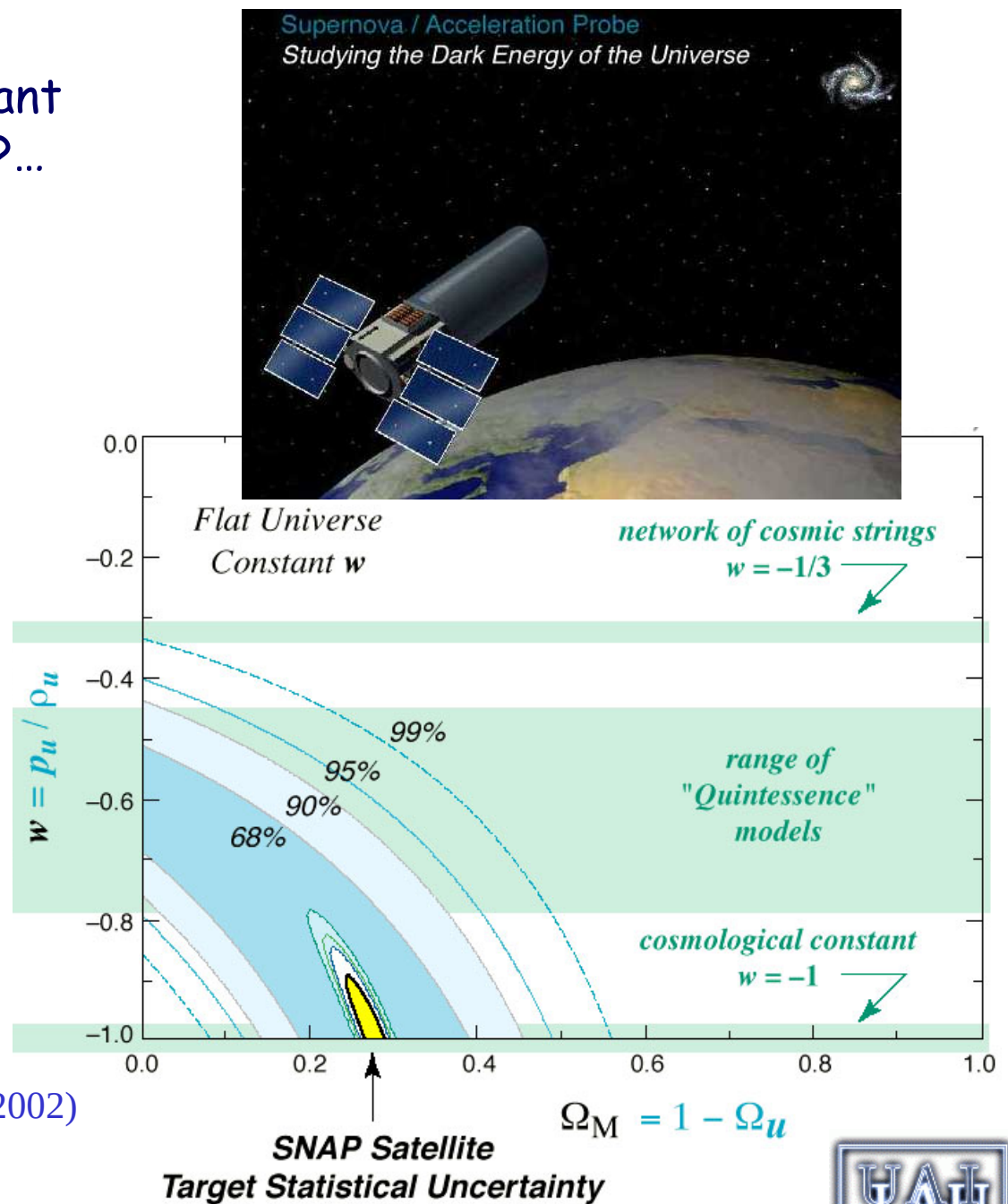
Adapted from Schmidt (2002)

Can we distinguish a constant Λ term from quintessence?

Not from current ground-based SN observations (combined with e.g. LSS)...

...or from future ground-based observations

Main goal of the SNAP satellite (launch ~2010?)



Adapted from Schmidt (2002)